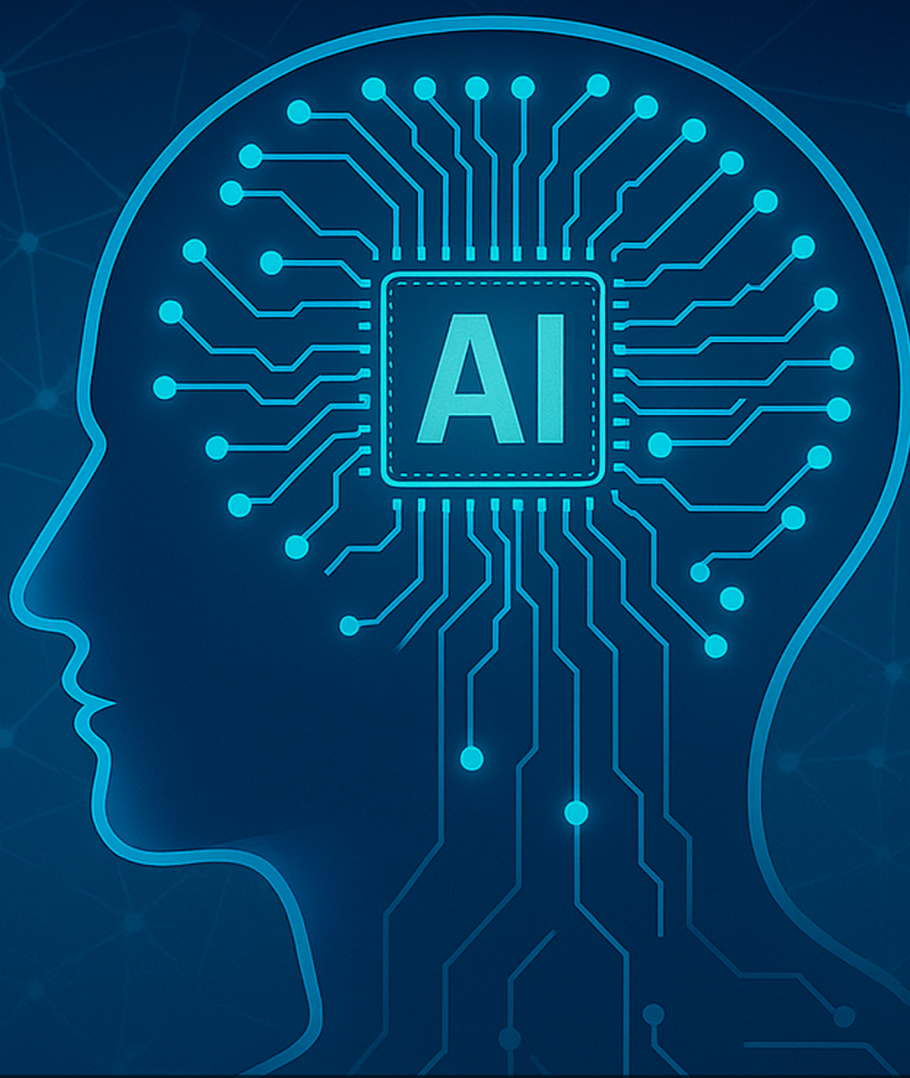


Sebastian Matysik, Krzysztof Jaworski

WYKORZYSTANIE AI W PRACY NAUKOWEJ



**Interdyscyplinarny
przewodnik dla doktorantów
i młodych naukowców**

Sebastian Matysik

Krzysztof Jaworski

Wykorzystanie AI w pracy naukowej

**Interdyscyplinarny przewodnik dla doktorantów i młodych
naukowców**

Sebastian Matysik

Krzysztof Jaworski

The Use of AI in Academic Work

An Interdisciplinary Guide for PhD Students and Young Researchers

Sebastian Matysik, Krzysztof Jaworski. Wykorzystanie AI w pracy naukowej. Interdyscyplinarny przewodnik dla doktorantów i młodych naukowców. Monograph. Szczecin: Scientific Publishing House (SPH), Centre of Sociological Research, 2025. - Bibliogr. – Illustr. –132 p.

Monografia powstała w wyniku współpracy autorów po warsztatach pt. „Wykorzystanie AI w pracy naukowej” współorganizowanych przez Szkołę Doktorską Uniwersytetu Szczecińskiego.



SZKOŁA DOKTORSKA
UNIwersytetu Szczecińskiego

Official reviewers of the monograph:

dr hab. Jarosław Korpysa, prof. US (Uniwersytet Szczeciński)

prof. dr hab. Maria Czerepaniak – Walczak (Uniwersytet Szczeciński)

All rights reserved. No part of this book may be reproduced or distributed in any form or by any means, or stored in a database or retrieval system, without the prior written permission of the publisher.

© 2025 Centre of Sociological Research

ISBN: 978-83-973513-6-3

DOI: 10.14254/978-83-973513-6-3/2025



Publishing House

Wydawnictwo Fundacji Centrum Badań Socjologicznych

Centre of Sociological Research

ul. Bolesława Śmiałego 22 lok. 27

70-347 Szczecin, Poland

e-mail: office@csr-pub.eu

<https://www.csr-pub.eu/>

Spis treści

Wstęp	11
Cel i motywacja napisania monografii	11
Metodyka opracowania.....	11
Część I Wprowadzenie teoretyczne	13
1. Wprowadzenie do AI w nauce	13
1.1. Definicje: AI, ML, LLM, generative AI.....	13
1.2. Modele językowe i sposoby ich szkolenia.....	15
1.3. Interfejsy narzędzi AI: chat, API, pluginy, rozszerzenia	17
1.4. Historia narzędzi wspomagających pisanie tekstów - od translatorów, przez Grammarly, po sztuczną inteligencję	20
1.5. Wzrost liczby publikacji naukowych napisanych z użyciem AI	21
2. Synergia AI i człowieka	24
2.1. Proces uczenia się	24
2.2. Uczenie się w perspektywie konstruktywizmu.....	25
2.3. Cykl uczenia się	25
2.4. Głębokie uczenie się – charakterystyka i praktyki wspierające rozwój	27
2.5. Iteracja, która łączy	28
2.6. Kompetencje w zakresie genAI	29
2.7. Kompetencje transwersalne	30
2.8. Model SAMR.....	32
2.9. Etyczny wymiar synergii	33
2.10. AI Act	34
2.11. Linie na piasku.....	34
3. Ograniczenia i zagrożenia wykorzystywania AI w pracy naukowej	36
3.1. Halucynacje i stronniczość	36
3.2. Skala błędnych cytowań generowanych przez AI	38
3.3. Zagadnienia dotyczące etyki.....	40
3.4. Prawa autorskie RODO/GDPR, odpowiedzialność autorów.....	41
4. Zmiany metodyczne z użyciem AI i przyszłe kierunki	43
4.1. Zmiana metodologii przeglądów literatury z wykorzystaniem AI	43
4.2. Przyszłe kierunki zmian związane z pisanie publikacji i recenzowaniem.....	44
5. Polityki wydawnictw naukowych i kwestie prawne	46

5.1. Stanowiska czołowych wydawnictw i czasopism naukowych	46
5.2. Kryteria autorstwa a użycie AI; deklaracje transparentności	53
5.3. Prawa autorskie do tekstu oraz modele licencjonowania	56
5.4. Dostosowanie narzędzi licencyjnych.....	57
Część II przewodnik praktyczny	59
6. Inżynieria promptów (prompt engineering)	59
6.1. Prompt engineering, a prompt learning	60
6.2. Struktura skutecznego procesu promptowania	61
6.3. Techniki prompt engineeringu.....	64
6.4. Sposoby budowania promptów.....	65
6.5. Wsparcie twórców genAI	66
6.6. Budowanie promptów w pracy naukowej i ograniczanie halucynacji.....	67
6.7. Dzielenie pracy na etapy w promptach.....	70
7. AI w pisaniu publikacji naukowych	73
7.1. Narzędzia AI wykorzystywane w pracy naukowej.....	74
7.2. Różne strategie współpracy z AI	84
7.3. Przykładowe ścieżki pracy z AI.....	85
7.4. Checklisty weryfikacji (fakty, referencje, plagiat)	86
8. Automatyzacja w pracy badawczej	89
8.1. Koncepcja „Sheets → AI → raport”.....	89
8.2. Klasyfikacja binarna abstraktów	90
8.3. Ekstrakcja informacji.....	91
8.4. Streszczenie robocze.....	92
8.5. Walidacja i replikowalność.....	92
8.6. Prywatność i zgodność.....	93
8.7. Najczęstsze problemy i rozwiązania.....	93
8.8. Ćwiczenie: mini-przegląd literatury	93
9. Detekcja treści generowanych przez AI	94
9.1. Detektory AI	94
9.2. Fałszywe alarmy i granice detekcji.....	98
10. Tworzenie i aktualizacja materiałów dydaktycznych.....	99
10.1. Tworzenie sylabusów i studiów przypadku.....	100
10.2. Tworzenie quizów, zadań i prezentacji.....	101

10.3. Metaprompty i refleksja nad uczeniem się	102
11. Najlepsze praktyki i rekomendacje dla młodych naukowców i doktorantów	104
11.1. Równowaga między efektywnością a rzetelnością.....	104
11.2. Matryca ryzyka/korzyści dla różnych etapów badań.....	106
11.3. Plan rozwoju kompetencji cyfrowych	108
Zakończenie	111
Spis tabel i rysunków	113
Literatura.....	115

Contents

Introduction	11
Purpose and motivation of writing the monograph	11
Methodology	11
Part I Theoretical introduction	13
1. Introduction to AI in science	13
1.1. Definitions: AI, ML, LLM, generative AI	13
1.2. Language models and training methods	15
1.3. AI tool interfaces: Chat, API, plugins, extensions	17
1.4. History of text-support tools – from translators, through Grammarly, to AI	20
1.5. Growth of scientific publications written with AI	21
2. Synergy of AI and humans	24
2.1. Learning process	24
2.2. Learning from the constructivist perspective	25
2.3. Learning cycle	25
2.4. Deep learning – characteristics and supporting practices	27
2.5. Iteration that connects	28
2.6. Competences in generative AI	29
2.7. Transversal competences	30
2.8. SAMR model	32
2.9. Ethical dimension of synergy	33
2.10. AI Act	34
2.11. Lines in the sand	34
3. Limitations and risks of using AI in scientific work	36
3.1. Hallucinations and bias	36
3.2. Scale of incorrect citations generated by AI	38
3.3. Ethical issues	40
3.4. Copyright, GDPR, and author responsibility	41
4. Methodological changes with AI and future directions	43
4.1. Changing literature review methodologies with AI	43
4.2 Future directions in writing and reviewing publications	44
5. Scientific publishing policies and legal issues	46
5.1. Positions of leading publishers and journals	46

5.2. Authorship criteria and use of AI; transparency declarations.....	53
5.3. Copyright of text and licensing models	56
5.4. Adaptation of licensing tools	57
Part II Practical guide	59
6. Prompt engineering.....	59
6.1. Prompt engineering a prompt learning	60
6.2. Structure of an effective prompting process	61
6.3. Prompt engineering techniques.....	64
6.4. Ways of building prompts.....	65
6.5. Supporting genAI creators	66
6.6. Building prompts in scientific work and reducing hallucinations	67
6.7. Dividing work into stages in prompts.....	70
7. AI in writing scientific publications.....	73
7.1. AI tools used in scientific work	74
7.2. Different strategies of collaboration with AI	84
7.3. Example pathways of working with AI	85
7.4. Verification checklists (facts, references, plagiarism).....	86
8. Automation in research work.....	89
8.1. Concept „Sheets → AI → report”	89
8.2. Binary classification of abstracts	90
8.3. Information extraction	91
8.4. Draft abstract.....	92
8.5. Validation and reproducibility	92
8.6. Privacy and compliance	93
8.7. Common problems and solutions	93
8.8. Exercise: mini literature review	93
9. Detection of AI-generated content.....	94
9.1. AI detectors.....	94
9.2. False alarms and detection limits.....	98
10. Creating and updating teaching materials	99
10.1. Creating syllabi and case studies	100
10.2. Creating quizzes, assignments, and presentations	101
10.3. Metaprompts and reflection on learning.....	102

11. Best practices and recommendations for young researchers and PhD students	104
11.1. Balance between efficiency and reliability	104
11.2. Risk/benefit matrix for different research stages	106
11.3. Digital competence development plan.....	108
Summary	111
List of tables and figures	113
References	115

Wstęp

Cel i motywacja napisania monografii

Celem monografii „Wykorzystanie AI w pracy naukowej” jest dostarczenie młodym badaczom i doktorantom spójnego, praktycznego i etycznie ugruntowanego przewodnika po zastosowaniach sztucznej inteligencji w całym cyklu badawczym od formułowania pytań i przeglądu literatury, przez pisanie i redakcję tekstu, po przygotowanie materiałów dydaktycznych i automatyzację zadań pomocniczych.

Impulsem do jej powstania jest gwałtowna adopcja narzędzi generatywnych, zwłaszcza dużych modeli językowych w pracy badawczej i redakcyjnej, czemu towarzyszą zagrożenia, takie jak halucynacje i stronniczość sztucznej inteligencji, a także nowe wymagania transparentności oraz stanowiska wydawnictw i redakcji w tym zakresie.

Monografia łączy zatem wyjaśnienia pojęć i trendów z procedurami, checklistami i mini zadaniami, aby ułatwić szybkie, odpowiedzialne wdrożenie AI w praktyce.

Kluczową motywacją jest także wyrównanie dostępu do wiedzy praktycznej. Wiele istniejących opracowań koncentruje się na technologiach, rzadziej pokazując jak bezpiecznie włączyć je do warsztatu naukowca i jak raportować ich użycie zgodnie z politykami czołowych wydawnictw. Monografia odpowiada na obie potrzeby: syntetyzuje podstawy teoretyczne i porządkuje najważniejsze decyzje metodologiczne oraz organizacyjne, które młody badacz podejmuje dziś, pracując z AI.

Pozycja powstała na bazie współpracy autorów publikacji w zakresie warsztatów realizowanych dla społeczności doktorantów i pracowników naukowych organizowanych w ramach Szkoły Doktorskiej Uniwersytetu Szczecińskiego. Zgromadzone doświadczenia i rozmowy z uczestnikami szkoleń pozwoliły na ustalenie struktury i treści monografii.

Sami autorzy są doktorantami szkoły doktorskiej, więc w pewnym sensie pozycja była pisana także z perspektywy własnych potrzeb oraz najbliższego otoczenia szkoły doktorskiej, w taki sposób, aby jak najlepiej odpowiedzieć na zachodzące zmiany i pojawiające się potrzeby.

Metodyka opracowania

Zastosowano połączenie scoping review i przeglądu narracyjnego artykułów, monografii opublikowanych w większości w latach 2019-2025 (stan na 15 sierpnia 2025 r.), obejmując badania nad AI/LLM, interakcję człowieka z AI, metodyką przeglądów oraz politykami wydawniczymi.

Część danych ilościowych dotyczy jawnie deklarowanego wykorzystania narzędzi AI w publikacjach, gdzie analizowano trzy źródła (Google Scholar, Web of Science, Scopus) przy użyciu spójnego zapytania obejmującego nazwy narzędzi (operator OR). Uwzględniono ograniczenia tych baz (szersze spektrum typów dokumentów w Google Scholar, bardziej rygorystyczne kryteria w WoS, Scopus).

Przeprowadzono przegląd stanowisk czołowych wydawnictw naukowych (Elsevier, Springer Nature, Taylor & Francis, Wiley, SAGE, IEEE, Emerald) oraz wybranych przykładów w Polsce. Ujęto wspólne minimum (brak autorstwa AI, obowiązek ujawnienia istotnego użycia, ostrożność wobec obrazów, danych, wymogi poufności w recenzji) i różnice w szczegółach, na przykład zakres ujawnień.

W procesie opracowania korzystano z narzędzi generatywnych jako wsparcia redakcyjno językowego i organizacyjnego, a także do wyszukiwania i syntezy literatury naukowej. Natomiast wszystkie aspekty merytoryczne, dane liczbowe i cytowania zostały zweryfikowane w źródłach pierwotnych, bibliografię i referencje sprawdzono ręcznie (DOI, tytuł, rocznik). Użycie AI jest jawne, zgodne z praktykami rekomendowanymi przez wydawców i wytyczne odpowiedzialnego stosowania AI w badaniach.

Stanowiska wydawnicze i parametry modeli dynamicznie się zmieniają, przedstawione porównania i zalecenia odzwierciedlają stan na kwiecień-sierpień 2025 r. Rekomendujemy weryfikację bieżących polityk konkretnego czasopisma i instytucji oraz dostosowanie procedur do specyfiki dyscypliny i danych.

Część I Wprowadzenie teoretyczne

Pierwsza część monografii koncentruje się na teoretycznych podstawach wykorzystania sztucznej inteligencji w nauce. Jej celem jest uporządkowanie kluczowych pojęć, przedstawienie najważniejszych nurtów rozwoju AI oraz ukazanie dynamiki ich przenikania do praktyki akademickiej. Zgromadzona tu wiedza pozwala lepiej zrozumieć zarówno potencjał, jak i ograniczenia nowych technologii, co stanowi fundament do dalszych rozważań o charakterze praktycznym. Dzięki temu młody badacz może świadomie osadzić własne doświadczenia w szerszym kontekście metodologicznym i etycznym.

1. Wprowadzenie do AI w nauce

Rozwój sztucznej inteligencji w ostatnich latach zmienił sposób prowadzenia badań naukowych, otwierając nowe możliwości w zakresie analizy danych, przeglądów literatury oraz pisania i redagowania tekstów. W szczególności generatywna AI oraz duże modele językowe stały się narzędziami, które coraz częściej wspierają pracę badaczy i doktorantów, zarówno w zadaniach koncepcyjnych, jak i w praktycznych aspektach przygotowywania publikacji. W niniejszym rozdziale przedstawiono definicje podstawowych pojęć związanych z AI, jej główne obszary zastosowań, a także dynamikę rosnącej obecności tych technologii w nauce.

1.1. Definicje: AI, ML, LLM, generative AI

Sztuczna inteligencja (AI)

Sztuczna inteligencja to obszar badań i praktyki inżynierskiej ukierunkowany na projektowanie systemów zdolnych wykonywać zadania, które zwykle wymagają ludzkiej inteligencji oraz sprawności takie jak postrzeganie, wnioskowanie, uczenie się, planowanie czy rozumienie języka. W literaturze ostatnich lat AI ujmuje się jako zestaw metod, w tym uczenie maszynowe i głębokie, które umożliwiają maszynom symulację funkcji poznawczych oraz adaptację do danych i kontekstu (Bajwa, Munir, Nori, Williams, 2021; Xu i in., 2021). Można zatem wskazać, że AI „symuluje ludzkie myślenie i zachowanie”, a praktycznie wykonuje zadania wymagające inteligencji, od rozpoznawania wzorców po rozumienie języka naturalnego (Bajwa i in., 2021; Xu i in., 2021).

Rozróżnia się wąską AI (systemy wyspecjalizowane do pojedynczych zadań) oraz postulat sztucznej inteligencji ogólnej (AGI) systemu zdolnego uczyć się szeroko i adaptować do nowych sytuacji w wielu obszarach. Współczesne definicje AGI akcentują formalne, obliczalne modele uczenia się i rozumowania, a także ścisłe powiązanie inteligencji z interakcją

ze środowiskiem (Bennett, 2022). Co ważne, badania nad AI coraz częściej podkreślają jej charakter socjotechniczny, AI funkcjonuje w sprzężeniu z praktykami społecznymi i instytucjonalnymi, co wpływa zarówno na sposób definiowania inteligencji maszynowej, jak i na wymagania etyczne i regulacyjne (Rezaev, 2021; Owczarczuk, 2023; Bennett, 2022).

Uczenie maszynowe (ML)

Uczenie maszynowe to podzbiór AI, w którym tworzy się modele uczące się na podstawie danych zamiast ręcznie programować reguły działania. W ujęciu współczesnym ML obejmuje uczenie nadzorowane (z etykietami), nienadzorowane (odkrywanie struktur), uczenie ze wzmocnieniem (uczenie polityki działania na podstawie nagród) oraz coraz powszechniejsze uczenie samonadzorowane, gdzie sygnał nadzorujący generowany jest z samych danych. Kluczowa jest zdolność do generalizacji, modele poprawiają swoje wyniki w miarę napływu danych i informacji zwrotnej (Barbierato, Gatti, 2024; Brnabic, Hess, 2021).

W najnowszych przeglądach uczenie maszynowe opisuje się jako projektowanie systemów, w których człowiek aktywnie koryguje i nadzoruje dane oraz działanie modelu – od etapu adnotacji po interwencje w trakcie uczenia. Takie podejście zwiększa jakość i wiarygodność modeli, szczególnie w zadaniach trudnych do pełnej automatyzacji. Aktualne syntezы porządkują terminologię, przedstawiają taksonomie technik i podkreślają, że w wielu zastosowaniach naukowych efektywność rośnie właśnie dzięki iteracjom człowiek-model (Mosqueira-Rey, Hernández-Pereira, Alonso-Ríos, Bobes-Bascarán, Fernández-Leal, 2023; Wu, Xiao, Sun, Zhang, Ma, He, 2022).

Duże modele językowe (LLM)

LLM (Large Language Models) to duże, wstępnie wytrenowane modele generatywne, zwykle oparte o architekturę transformera, uczące się przewidywania kolejnych tokenów w tekście na podstawie bardzo rozległych zbiorów. Token to podstawowa jednostka tekstu, na której operują duże modele językowe może odpowiadać całemu słowu, jego części lub nawet pojedynczemu znakowi. Modele uczą się przewidywania kolejnych tokenów w sekwencji, dlatego liczba tokenów wyznacza maksymalny zakres tekstu, jaki mogą przetworzyć jednorazowo. Współczesne przeglądy definiują LLM jako klasę modeli z reguły liczonych w dziesiątkach miliardów parametrów, które dzięki skalowaniu danych i obliczeń ujawniają zdolności uogólnione od streszczania i odpowiadania na pytania, przez tłumaczenia, po programowanie (Zhao i in., 2023; Minaee, Mikolov, Nikzad, Chenaghlu, Socher, Amatriain, Gao, 2024). W zastosowaniach naukowych LLM postrzega się jako szczególny przypadek AI

generatywnej ukierunkowanej na język naturalny (Lu, Peng, Cohen, Ghassemi, Weng, Tian, 2024).

Kluczową rolę odgrywa skalowanie, ponieważ między innymi prace nad GPT3 pokazały, że powiększanie modelu i danych poprawia wyniki w trybie „fewshot” bez dodatkowego dostrajania (Brown i in., 2020). Nowsze analizy praw skali doprecyzowują, że dla danego budżetu obliczeń istnieje relacja między rozmiarem modelu a liczbą tokenów potrzebnych do optymalnego treningu, tak zwana „Chinchilla scaling law” (Hoffmann i in., 2022). Równolegle, raport o modelach fundamentalnych opisuje LLM jako modele bazowe trenowane na szerokich danych, łatwe do adaptacji do wielu zadań, ale też niosące ryzyko ujednolicenia błędów w zastosowaniach praktycznych (Bommasani i in., 2021).

Generatywna AI (generative AI)

Generatywna AI obejmuje techniki obliczeniowe, które potrafią tworzyć nowe treści (tekst, obrazy, dźwięk, wideo czy struktury molekularne) na podstawie rozkładów wyuczonych z danych. Najnowsze opracowania konceptualizują GenAI jako element systemów socjotechnicznych i wskazują jej gwałtowną dyfuzję od DALL-E 2, przez GPT-4, po asystentów kodowania, co zmienia praktyki pracy i komunikacji (Feuerriegel, Hartmann, Janiesch, Zschech, 2024). Równolegle ukazały się przeglądy historyczne pokazujące etapy rozwoju GenAI na przestrzeni siedmiu dekad oraz powiązania z modelami fundamentalnymi (He, Cao, Tan, 2025).

Do najważniejszych paradygmatów generatywnych należą dziś między innymi modele dyfuzyjne które osiągnęły przełomową jakość syntezy (Ho, Jain, Abbeel, 2020) oraz szerokie spektrum metod opisanych w aktualnych przeglądach, na przykład mechanizmy przyspieszania próbkowania czy łączenia z innymi klasami modeli (Yang i in., 2022). W praktyce to właśnie modele dyfuzyjne i wielkoskalowe LLM, często multimodalne napędzają współczesny krajobraz GenAI (Ho, Jain, Abbeel, 2020; Yang i in., 2022).

Szybka adopcja GenAI pociąga za sobą pytania o wiarygodność, stronniczość, przejrzystość i ramy regulacyjne w nauce i poza nią. Nowsze opracowania przeglądowe rekomendują budowę pewnych ram regulacyjnych i praktyk odpowiedzialnego wdrażania na przykład w medycynie, administracji, edukacji, tak aby maksymalizować korzyści i jednocześnie minimalizować ryzyka społeczne (Owczarczuk, 2023).

1.2. Modele językowe i sposoby ich szkolenia

Duże modele językowe (LLM) oparte na architekturze transformera stały się uniwersalnym narzędziem badawczym, ponieważ w pewnym sensie potrafią rozumieć i generować tekst, a po niewielkiej adaptacji wspierają zadania od streszczania literatury po tworzenie prototypów analiz (Minaee i in., 2024; Zhao i in., 2023).

Transformery wykorzystują samouwagę (self-attention), która pozwala modelowi przeanalizować zobaczyć dowolny fragment zdania i uczyć się zależności w długim kontekście. W praktyce dało to dużą przewagę nad starszymi podejściami takimi jak RNN, czy modele statystyczne, które gorzej skalują się na długie teksty (Minaee i in., 2024; Zhao i in., 2023). Wersje otwarte na przykład LLaMA, Llama 2 umożliwiły szybkie eksperymenty akademickie (Touvron i in., 2023a, 2023b). Z kolei modele wielkoskalowe, takie jak PaLM, pokazały, że wzrost danych i mocy obliczeniowej przekłada się na wyraźny skok jakości (Chowdhery i in., 2022).

Rodziny modeli:

- Encoderonly (na przykład BERT, DeBERTa) - najlepsze, gdy celem jest rozumienie, klasyfikacja, ekstrakcja informacji. Uczą się zwykle przez maskowanie tokenów (He i in., 2021).
- Decoderonly (na przykład GPT3/4) - przeznaczone do generacji, trenują się przez przewidywanie kolejnego tokenu i dobrze działają w trybie zero/fewshot (Brown i in., 2020; Minaee i in., 2024).
- Encoderdecoder (na przykład T5/BART) - wygodne w transdukcji tekstu do tekstu, czyli tłumaczenie, streszczanie, często wybierane w zadaniach wymagających zarówno rozumienia, jak i generowania (Zhao i in., 2023).

Sposoby szkolenia modeli:

- 1) Pretrenowanie samonadzorowane - model uczy się na nieoznaczonych danych, albo maskuje część wyrazów (warianty typu BERT, DeBERTa), albo przewiduje następny token (warianty GPT). Duża skala danych i obliczeń jest kluczowa dla jakości (He i in., 2021; Chowdhery i in., 2022; Minaee i in., 2024).
- 2) Adaptacja do zadań - bez pełnego przeuczenia
 - Promptowanie i uczenie kontekstowe: wystarczy dobra treść polecenia, przykłady w wejściu (Brown i in., 2020).
 - Instruction tuning: dostrajanie na zbiorach poleceń i odpowiedzi zwiększa podporządkowanie modelu poleceniom (Wei i in., 2022; Chung i in., 2024).

- RLHF / DPO: porządkowanie zachowania pod preferencje ludzkie albo z modelem nagrody (RLHF), albo bezpośrednio z par preferencji (DPO) (Ouyang i in., 2022; Rafailov, Sharma, Mitchell, Ermon, Manning, Finn, 2023).
- PEFT - efektywna parametryczna adaptacja: zamiast ruszać wszystkie wagi, modyfikujemy małe macierze (LoRA) lub dostrajamy model w niskiej precyzji (QLoRA) to znacznie obniża koszty sprzętowe (Hu i in., 2022; Dettmers i in., 2023).
- Wzmacnianie wiedzy przez wyszukiwanie (RAG, RETRO): model w trakcie działania dociąga fragmenty tekstów z zewnętrznej bazy, co poprawia aktualność i faktualność bez implementacji wszystkiego w parametry (Borgeaud i in., 2022).

Główne wyzwania związane z trenowaniem modeli:

- Koszt danych i obliczeń - skuteczność rośnie wraz ze skalą, ale budżet jest ograniczony. Zasada „Chinchilla” sugeruje, by dla danego budżetu intensywniej zwiększać liczbę tokenów niż wyłącznie rozmiar modelu (Hoffmann i in., 2022). PEFT (LoRA/QLoRA) pomaga w praktyce (Hu i in., 2022; Dettmers i in., 2023).
- Długi kontekst - klasyczna uwaga ma koszt kwadratowy względem długości. Ulepszenia, takie jak FlashAttention, czyli szybsze i oszczędniejsze liczenie uwagi, ALiBi, czyli ekstrapolacja do dłuższych sekwencji czy architektury longcontext (LongNet), zwiększają użyteczny kontekst i wydajność (Dao i in., 2022; Press, Smith, Lewis, 2022; Ding i in., 2023).
- Adaptacja do różnych obszarów - dziedzinach specjalistycznych na przykład w przypadku medycyny, prawa najlepiej łączyć małe dostrajanie parametrów (PEFT) z zewnętrznym wyszukiwaniem źródeł (RETRO, RAG), co poprawia trafność i redukuje halucynacje (Borgeaud i in., 2022; Minaee i in., 2024).

1.3. Interfejsy narzędzi AI: chat, API, pluginy, rozszerzenia

W ostatnich latach duże modele językowe (LLM) udostępniane poprzez API zdobyły ogromną popularność w nauce i technice, znajdując szerokie zastosowanie w różnych dyscyplinach (Zhang i in., 2025). Coraz więcej badań opisuje integrację LLM takich jak GPT-3, GPT-4 z narzędziami badawczymi i workflow akademickimi (Wang, Hu, Huang, Li, Zhang, Ning, Zhu, Li, Ye, 2024).

Systematyczne przeglądy literatury wskazują, że modele te wykorzystywane są w różnych dziedzinach na przykład geoinformatyka, chemia, nauki biologiczne do automatyzacji analiz, generowania tekstu naukowego czy wspomagania eksperymentów. Modele te stają się elementem infrastruktury naukowej, wspierając naukowców na wielu etapach procesu badawczego. (Wang i in., 2024; Zhang i in., 2025).

LLM stają się elementem infrastruktury badawczej: automatyzują analizy, wspierają generowanie i redakcję tekstu naukowego oraz asystują w prowadzeniu eksperymentów (Wang i in., 2024; Zhang i in., 2025). Interfejs czatowy umożliwia iteracyjną pracę w języku naturalnym, czyli konsultacje, doprecyzowanie, podsumowania. API, czyli interfejs, który pozwala programom wymieniać dane i współpracować ze sobą, daje możliwość na integracji programistycznej i automatyzacji przetwarzania na dużą skalę. Pluginy zaś dają możliwość sięgania po zewnętrzne narzędzia i aktualne dane w trakcie interakcji, a rozszerzenia wbudowanie modeli bezpośrednio w aplikacje na przykład edytory, IDE, przeglądarki, bez zmiany środowiska pracy (Wang i in., 2024).

Liczne prace empiryczne wskazują, że LLM mogą pełnić rolę scientific copilots: streszczają treści publikacji, wydobywają informacje, generują i tłumaczą fragmenty tekstu, wspomagają programowanie na przykład przez szkielety analiz, a nawet planują kroki badawcze (Zhang i in., 2025). Integracja z systemami zarządzania danymi i oprogramowaniem statystycznym usprawnia przepływ pracy i pozwala skupić się na interpretacji wyników, przy zachowaniu krytycznej weryfikacji treści generowanych przez modele (Wang i in., 2024; Zhang i in., 2025).

Interfejs czatowy sprzyja iteracyjnej pracy w języku naturalnym doprecyzowania, prompt chaining, wieloetapowe wyjaśnienia i, co potwierdza Jakesch i in. (2023), może podnosić produktywność oraz pewność siebie piszących, ale zarazem zmienia dynamikę procesu pisarskiego, potrzebę projektowania wsparcia AI tak, by wzmacniało sprawczość użytkownika. Badania Li, Liang, Peng, Yin (2024) pokazują, że asysta generatywnej AI poprawia wydajność i samoocenę piszących, a sam projekt narzędzia ma znaczenie dla jakości efektu. Jednocześnie prace empiryczne ujawniają, że współpisanie z modelem może kierunkować treść i opinie autorów, na przykład przez ryzyko wpływu asystenta w dialogu, co wzmacnia potrzebę transparentnych interfejsów i praktyk weryfikacji (Jakesch i in., 2023; Jiang i in., 2024).

W ujęciu technicznym API umożliwia trzy główne wzorce:

- (1) RAG (Retrieval-Augmented Generation) - dołączanie zewnętrznej wiedzy przez wyszukiwanie, indeks (ogranicza halucynacje, pozwala aktualizować model danymi domenowymi),
- (2) tool-use wywołania funkcji - pozwalające modelowi w trakcie generacji korzystać z narzędzi i usług (kalkulator, wyszukiwarka, bazy),
- (3) architektury agentowe - łączenie rozumowania i działania w pętlach decyzyjnych.

Ostatnie przeglądy potwierdzają, że RAG stał się standardem inżynierskim dla aplikacji naukowych, bo systematycznie poprawia trafność i uwiarygadnia odpowiedzi (Gao i in., 2023; Fan i in., 2024). Z kolei tooluse i planowanie akcji opisują prace Toolformer - samouczenie się wywołań API przez LLM oraz ReAct - przeplatanie śladów rozumowania i działań, stanowiąc fundament do samodzielności LLM w nauce (Schick i in., 2023; Yao i in., 2023).

Konkretny dowód z laboratorium: system Coscientist (Nature) integruje GPT4 z wyszukiwaniem, wykonywaniem kodu i API do sterowania aparaturą, by autonomicznie planować i przeprowadzać reakcje chemiczne; podobnie LLM wygenerują skrypty dla robotów laboratoryjnych na podstawie instrukcji w języku naturalnym (Boiko i in., 2023; Inagaki i in., 2023).

Pluginy (wtyczki) rozszerzają możliwości modeli, podczas dialogu wywołuje usługi stron trzecich, czyli API baz publikacji, obliczenia, przegląd sieci, a wynik włącza do odpowiedzi.

Pierwsze przekrojowe badanie bezpieczeństwa ekosystemu wtyczek ChatGPT (ASE 2024) ujawniło luki w uwierzalnianiu i ochronie danych oraz nierówny rozkład funkcjonalności wtyczek; autorzy proponują model oceny bezpieczeństwa na trzech warstwach (użytkownik-twórca-operator sklepu) i zalecenia dotyczące „least privilege” oraz przeglądów bezpieczeństwa (Yan i in., 2024).

W praktyce naukowej oznacza to, że korzystając z wtyczek na przykład do wyszukiwania literatury i pobierania pełnych tekstów, instytucje powinny egzekwować kontrolę uprawnień i audyt logów wywołań, na przykład zgodność z politykami danych.

Rozszerzenia osadzają LLM bezpośrednio w narzędziach pracy, czyli w edytorach, IDE, przeglądarkach, co redukuje tarcie poznawcze i przyspiesza przepływ pracy. W IDE najsilniejszym empirycznym dowodem są badania kontrolne GitHub Copilot: programiści z dostępem do asystenta kończyli zadania szybciej, w jednym z RCT średnio ~56% szybciej, a kolejne eksperymenty terenowe i RCT w środowisku pracy potwierdzają wzrost produktywności i subiektywnej satysfakcji (Peng i in., 2023; RCT 2024).

W pakietach biurowych na przykład Microsoft 365 Copilot architektura łączy kontekst użytkownika z modelem w obrębie granic usług M365, z kontrolą dostępu na poziomie tożsamości i uprawnień, co jest kluczowe przy pracy na dokumentach i danych badawczych (Microsoft, 2025).

1.4. Historia narzędzi wspomagających pisanie tekstów - od translatorów, przez Grammarly, po sztuczną inteligencję

Współczesny ekosystem narzędzi można ująć w cztery kategorie:

- translatory maszynowe,
- korektory gramatycznostylistyczne na przykład Grammarly,
- parafrazery i inteligentne sugestie
- generatywne modele językowe (LLM).

Badania w dydaktyce języków i pisaniu akademickim pokazują przesunięcie akcentu od poprawiania istniejącego tekstu do proaktywnego współtworzenia treści i stylu (GodwinJones, 2022; Gustilo, Ong, Lapinid, 2024).

Upowszechnienie neuronowego tłumaczenia około 2016 r. podniosło płynność i adekwatność przekładów, co w praktyce otworzyło drogę do pisania w swoim języku i postedycji wersji angielskiej. Empiryczne badania w szkołach średnich wykazały, że użycie Google Translate z postedycją istotnie podnosi miary jakości pisania w L2 (Cancino, Panes, 2021). Dla autorów spoza kręgu anglojęzycznego narzędzie to stało się rutynowym wsparciem przygotowania szkiców i rozumienia literatury (Cancino, Panes, 2021).

W porównaniu z klasycznym sprawdzaniem pisowni, nowa generacja narzędzi dostarcza bieżących, kontekstowych podpowiedzi od zgody podmiotu z orzeczeniem po skracanie zdań i dopasowanie tonu wypowiedzi. Badanie Fitria (2021) wskazuje dużą poprawę ocenianej jakości prac studentów po zastosowaniu sugestii Grammarly, co ilustruje potencjał takich narzędzi w eliminowaniu błędów powierzchniowych i porządkowaniu stylu. Równolegle literatura opisuje integrację podpowiadania i autouzupełniania w szerszych scenariuszach nauczania pisania (GodwinJones, 2022).

Narzędzia parafrazujące oraz funkcje autokompletacji przesuwają ciężar z korekty na aktywną redakcję, co z jednej strony sprzyja klarowności tekstu, z drugiej rodzi pytania o integralność wkładu autorskiego. Przegląd polityk i praktyk uczelni wskazuje, że środowisko akademickie postrzega te rozwiązania jako użyteczne, ale wymagające transparentności i wytycznych stosowania (Gustilo i in., 2024).

Modele LLM na przykład GPT4, GPT-5 umożliwiają tworzenie streszczeń, przeglądów literatury czy wstępnych szkiców sekcji metod i dyskusji, co znacząco skraca się czas wstępnej redakcji. Jednocześnie literatura podkreśla ryzyka: halucynacje, niezamierzony autoplgiat oraz ujednolicanie stylu, stąd potrzeba weryfikacji faktów i jasnego ujawniania użycia narzędzi (Dehouche, 2021; Gustilo i in., 2024).

Systematyczne przeglądy dokumentują szybkie włączanie LLM do praktyk badawczych w wielu obszarach, w tym automatyzację analiz, wspomaganie pisanie i syntezę wyników (Wang, Hu, Huang, Li, Zhang, Ning, Zhu, Li, & Ye, 2024). Dla doktorantów także oznacza to przejście od narzędzi korekty do asystentów procesu pisarskiego pod warunkiem zachowania nadzoru merytorycznego człowieka, postędy i jawności (Wang i in., 2024; Gustilo i in., 2024).

W pracy nad manuskryptem warto:

- (a) łączyć translatory NMT (*Neural Machine Translation*) z uważną postędy (Cancino & Panes, 2021),
- (b) używać korektorów do usprawniania stylu i spójności językowej,
- (c) wykorzystać LLM do szkicowania lub przeformułowania fragmentów, ale każdą treść weryfikować i ujawniać użycie narzędzi w sekcji metodyka, podziękowania (Gustilo i in., 2024; Dehouche, 2021).

1.5. Wzrost liczby publikacji naukowych napisanych z użyciem AI

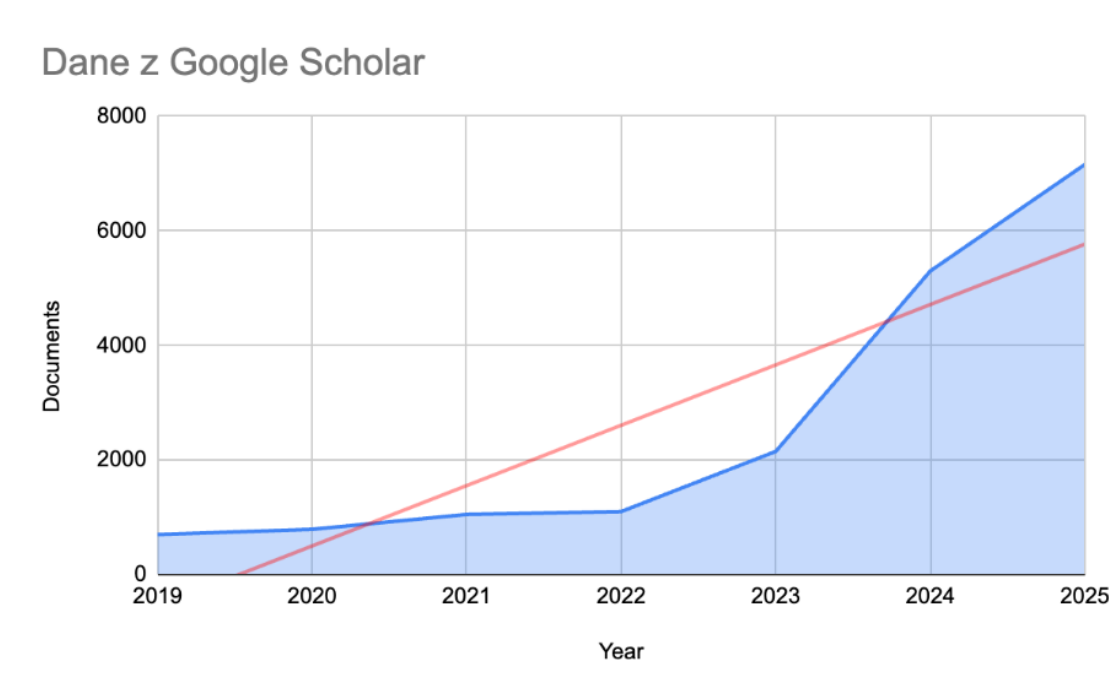
Dokonano analizy, jak szybko rośnie liczba prac naukowych, w których autorzy jawnie deklarują użycie narzędzi AI w procesie badawczym lub redakcyjnym. Analiza obejmuje trzy źródła: Google Scholar, Web of Science oraz Scopus i dotyczy lat 2019-2025, przy czym rok 2025 liczony jest do 15 sierpnia. We wszystkich bazach zastosowano to samo zapytanie z operatorem OR obejmujące dokładne nazwy narzędzi:

"scite.ai" OR "notebooklm" OR "SciSpace" OR "researchrabbitapp" OR "thesify.ai" OR "paperpal.com" OR "paperguide.ai" OR "elicit.com" OR "chatpdf.com" OR "textero.io" OR "consensus.app" OR "epsilon-ai" OR "jenni.ai" OR "getcoralai.com" OR "scholarai.io" OR "powerdrill.ai"

Nie dokonano zatem pomiaru liczby publikacji naukowych generowanych przez AI, natomiast prace, które wprost wymieniają konkretne narzędzia w metodach, podziękowaniach lub oświadczeniach.

Z danych z Google Scholar, które przedstawiono na rysunku 1.1 odnotowano kolejno: 699 publikacji w 2019 roku, 790 w 2020 roku (wzrost o 13,02 procent rok do roku), 1 050 w 2021 roku (wzrost o 32,91 procent), 1 100 w 2022 roku (wzrost o 4,76 procent), 2 150 w 2023 roku (wzrost o 95,45 procent), 5 300 w 2024 roku (wzrost o 146,51 procent) oraz 7 160 w 2025 roku do dnia piętnastego sierpnia (wzrost o 35,09 procent względem 2024 roku). Stosunek wartości z 2024 roku do wartości z 2019 roku wynosi 7,582, a stosunek wartości z 2025 roku

do 2019 roku 10,243. Średnie roczne tempo wzrostu w latach 2019-2024 wyniosło 49,95 procent, co potwierdza przyspieszający charakter zjawiska po 2022 roku.



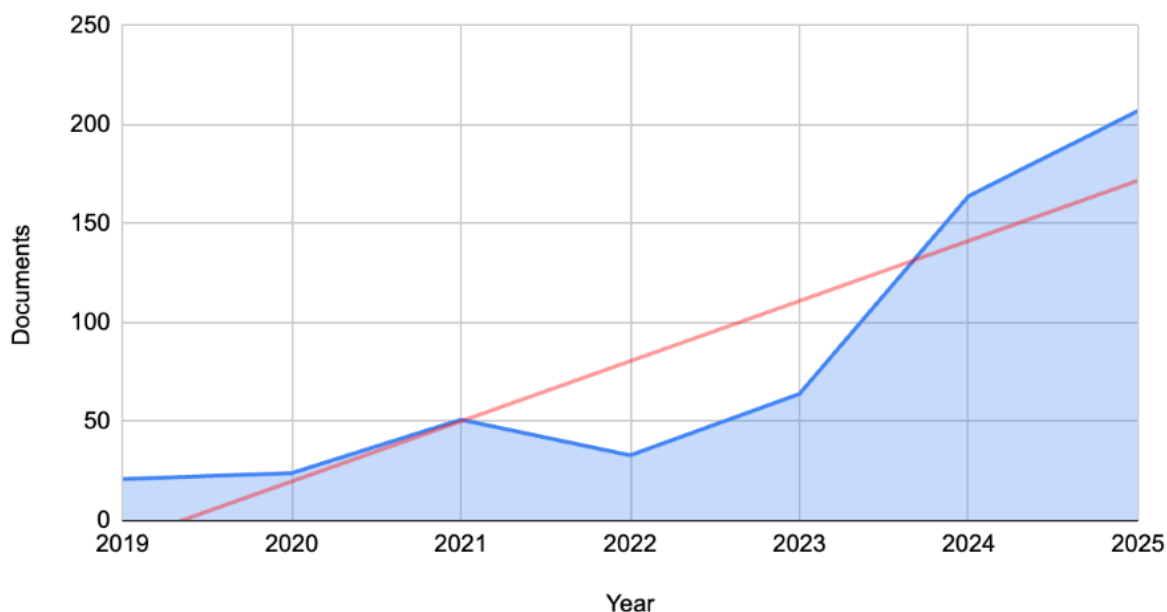
Rys. 1.1. Liczba publikacji wymieniających narzędzia AI w bazie Google Scholar (2019-2025)

Źródło: opracowanie własne na podstawie danych z Google Scholar (data dostępu: 15.08.2025)

Dane z Web of Science i Scopus potwierdzają tę dynamikę, choć na niższym poziomie bezwzględnym, co przedstawiono na rysunku 1.2. Odnotowano 21 publikacji w 2019 roku, 24 w 2020 roku (wzrost o 14,29 procent), 51 w 2021 roku (wzrost o 112,50 procent), 33 w 2022 roku (spadek o 35,29 procent), 64 w 2023 roku (wzrost o 93,94 procent), 164 w 2024 roku (wzrost o 156,25 procent) oraz 207 w 2025 roku do dnia piętnastego sierpnia (wzrost o 26,22 procent względem 2024 roku).

Stosunek wartości z 2024 roku do 2019 roku wynosi 7,810, a stosunek wartości z 2025 roku do 2019 roku 9,857. Średnie roczne tempo wzrostu w latach 2019–2024 wyniosło 50,84 procent. Jedyne spadki w 2022 roku wskazuje na przejściowe zahamowanie lub opóźnienia indeksacyjne, po którym nastąpił wyraźny skok w kolejnych latach.

Dane z Web of Science i Scopus



Rys. 1.2. Liczba publikacji wymieniających narzędzia AI w bazach Web of Science i Scopus (2019-2025)

Źródło: opracowanie własne na podstawie danych z Web of Science i Scopus (data dostępu: 15.08.2025)

Zbieżność trendów między bazami sugeruje, że rok 2023 był punktem zwrotnym: narzędzia do analizy literatury, streszczania i wsparcia redakcyjnego zaczęły być systematycznie ujawniane w publikacjach. W praktyce badawczej pożądane jest konsekwentne wskazywanie nazw użytych narzędzi wraz z datą dostępu lub wersją, dokumentowanie zapytań i parametrów oraz weryfikowanie treści generowanych przez systemy AI względem źródeł pierwotnych. Dobór narzędzia powinien odpowiadać specyfice zadania, niezależnie od tego, czy dotyczy wyszukiwania literatury, redakcji tekstu, czy organizacji wiedzy.

Należy podkreślić ograniczenia pomiaru: zliczane są wyłącznie publikacje jawnie wskazujące nazwę narzędzia (część użycia pozostaje niewidoczna lub zapisana inaczej), Google Scholar obejmuje szersze spektrum typów dokumentów, co sprzyja przeszacowaniu, natomiast WoS, Scopus są bardziej rygorystyczne. Mimo tych różnic, zebrane dane prowadzą do jednoznacznego wniosku: w latach 2019-2025 utrwała się dynamiczny, przyspieszający wzrost deklarowanego wykorzystania narzędzi AI w procesie publikacyjnym, a wartości odnotowane do 15 sierpnia 2025 r. przewyższają poziomy z 2024 r. w obu zestawach danych.

Podsumowanie rozdziału 1

Rozdział ten ukazał, że sztuczna inteligencja, a zwłaszcza modele generatywne, stają się integralną częścią współczesnego warsztatu naukowego. Wyjaśniono podstawowe pojęcia i przedstawiono przykłady ich zastosowań w praktyce badawczej, podkreślając zarówno korzyści, jak i ograniczenia związane z tymi technologiami. Zebrane analizy wskazują, że dynamiczny wzrost liczby publikacji opartych na AI dowodzi trwałej zmiany w sposobie tworzenia wiedzy, która wymaga od młodych naukowców umiejętności krytycznej oceny, przejrzystości i odpowiedzialności. AI w nauce nie jest już jedynie narzędziem pomocniczym, lecz staje się pełnoprawnym elementem procesu badawczego, współkształtując jego metodologię i praktykę.

2. Synergia AI i człowieka

Technologia zmienia się i rozwija w wykładniczym tempie. Ciągłe zmiany interfejsów i nowe funkcjonalności są czymś, co może przytłaczać. W pracy badaczy i dydaktyków potrzebne są ponadczasowe ramy, które są uniwersalne. Takie podejście może pozwolić wzmacniać obszar świadomego sięgania po cyfrowe rozwiązania. Pomimo dynamiki zmian warto osadzić korzystanie z generatywnej sztucznej inteligencji na solidnym fundamencie. Widoczna jest potrzeba znalezienia takiej perspektywy, która daje możliwość podejmowania decyzji i działań niezależnie od tego jak szybko i w jakim zakresie będzie dokonywał się rozwój narzędzi opartych na genAI.

Próbie znalezienia tego fundamentu, na którym można budować synergię możliwości człowieka i maszyny rozpoczyna się od założenia, że doktoranci, młodzi naukowcy są zanurzeni w ustawiczny proces uczenia się. To jest aspekt, który umożliwia, że podejmowane projekty badawcze osiągają zamierzone cele, a ich realizatorzy rozwijają zasoby kompetencji. Korzystanie z generatywnej sztucznej inteligencji jest obszarem, który może te działania wzmacniać i pogłębiać. Może też niekorzystnie wpływać na warsztat dydaktyczny i naukowy.

2.1. Proces uczenia się

Do tej pory opisane zostały definicje i odpowiedzi na pytanie czym jest sztuczna inteligencja w sensie technicznym. Jednak jako technologia przenikająca coraz częściej proces realizacji badań naukowych wymaga szerszego spojrzenia. „Generatywna AI nie jest neutralnym narzędziem; jest siłą kształtującą naturę produkcji i rozpowszechniania wiedzy. Rozumiejąc jej mechanikę, wykorzystując jej mocne strony i łagodząc jej słabości, badacze

mogą wykorzystać potencjał AI, zachowując jednocześnie integralność dociekań akademickich" (Jemieliński i in., 2025, s. 10) Istotnym elementem, który dotyczy korzystania z różnych technologii AI jest rozpoznawanie, odświeżanie i poszerzanie, jaki jest nasz ludzki potencjał. Jedno i drugie, czyli pogłębianie wiedzy na temat technologii i nas samych jest istotnym elementem umożliwiającym współpracę z AI opartą na synergii.

Jednocześnie kluczowe jest założenie, że jako młodzi naukowcy, badacze, dydaktycy bierzemy udział w procesie ustawicznego uczenia się. Ten proces przenika naszą pracę koncepcyjną, badawczą, upowszechniającą, ewaluacyjną i relacje międzyludzkie. Dlatego uczenie się jest w tym rozdziale jednym z obszarów na którym warto się zatrzymać. Zostanie odczytane w duchu konstrukcjonizmu. Następnie zostaną przeanalizowane kompetencje transversalne, aby pokazać jak różne poziomy korzystania z nowych technologii angażują różne poziomy kompetencji i jakie warunki należy spełnić, aby proces synergii był oparty na świadomym korzystaniu z zasobów generatywnej sztucznej inteligencji i ludzkiego potencjału. Poniższym rozważaniom towarzyszyć będzie głos Seymoura Paperta: „prawdziwa edukacja komputerowa to nie jest po prostu wiedza, jak stosować komputery i pomysły komputerowe. Jest to wiedza o tym, kiedy należy je stosować” (Papert, 1981, s. 176).

2.2. Uczenie się w perspektywie konstruktywizmu

Uczenie się jest aktywnym procesem, w ramach którego ludzie konstruują swoją wiedzę na podstawie doświadczeń w świecie. (Bruckman, Resnick, s. 214). Spojrzenie to wywodzi się z konstruktywizmu piagetowskiego. Zostało rozwinięte przez Seymoura Paperta, który dodał drugi, ważny dla nas aspekt, mówiący, że ludzie konstruują nową wiedzę szczególnie efektywnie wtedy, kiedy angażują się w konstruowanie artefaktów, które mają dla nich osobiste znaczenie (Bruckman, Resnick, s. 214).

Po pierwsze, w takim ujęciu uczenie się wymaga aktywnej postawy osoby, która w celowy sposób przeobraża własną sieć poznawczą. Po drugie, proces ten jest sprzężony z zewnętrznymi wytworami na przykład przyjmować postać eseju, analizy i syntezy wyników badań, programu komputerowego czy innych rozwiązań. Artefaktem mogą być też prompty, czy inne opracowane przez nas narzędzia i metody, które wspierają pracę z generatywną sztuczną inteligencją.

2.3. Cykl uczenia się

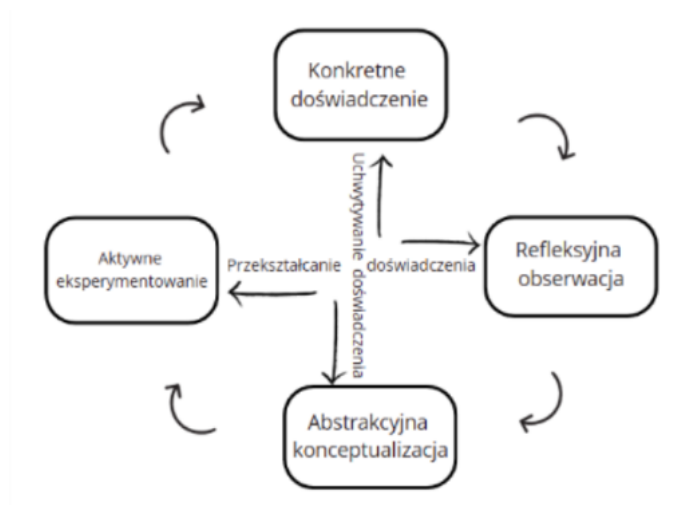
Proces uczenia się w przyjętym duchu często obrazowany jest jako cykl, w ramach którego przebiegają jego kolejne etapy. Mitchel Resnick zaproponował „spirale kreatywnego uczenia się”, co przedstawiono na rysunku 2.1.



Rys. 2.1. Spirala kreatywnego uczenia się

Źródło: (Resnick, 2017)

Jest to proces, który rozpoczyna się od idei uchwyconej w głowie, która następnie materializuje się poprzez działanie, zabawę, dzielenie z innymi, refleksję, aby następnie znów rozpocząć kolejny pogłębiony krąg. Schemat ten, który przedstawiono na rysunku 2.2 w sensie modelu znajduje się bardzo blisko cyklu uczenia się opartego na doświadczeniu, zaproponowanego przez Davida Kolba.



Rys. 2.2. Cykl Uczenia się opartego na doświadczeniu

Źródło: (Kolb, 2024)

Oba te schematy podkreślają wartość osobistego doświadczania oraz pogłębionej refleksji i (re)interpretacji. Opisują w sposób uniwersalny ustawiczny charakter uczenia się oraz wartość głębokiego zanurzenia się w działalność poznawczą i eksploracyjną.

Schemat następujących po sobie etapów wskazuje, że uczenie się jest oparte na powtórzeniach, w ramach których pogłębiamy kompetencje. Dbanie o swój własny warsztat pracy wymaga wysiłku, zaangażowania w proces myślenia, a także jak mawiał Seymour Papert „myślenia o myśleniu” (Papert, s. 39), co pozwala w krytyczny i kreatywny sposób budować

sieć poznawczą i operacyjną, którą wykorzystujemy między innymi w pracy badawczej. Dlatego ważne jest świadome osadzenie wartości pogłębionej pracy, co pozwala rozwijać i doskonalić w sobie zdolność głębokiego myślenia, oceny i refleksji. "Badacze muszą zachować rygorystyczne standardy analityczne, traktując wyniki generowane przez AI jako wstępne spostrzeżenia wymagające dokładnego zbadania, walidacji i dopracowania. Celem jest wzmocnienie, a nie zlecenie na zewnątrz, procesu intelektualnego, zapewniając, że ludzka ekspertyza, wiedza dziedzinowa i osąd akademicki kierują konceptualizacją badań". (Jemieliński, i in., 2025, s. 30)

Jednocześnie, wykorzystanie generatywnej AI niesie ze sobą ryzyka poznawcze i etyczne. Jednym z nich jest zjawisko *cognitive offloading*, rozumiane jest jako korzystanie z zewnętrznych środków w celu realizacji zadań poznawczych (Jose, Cherian, Verghis, 2025). Może to prowadzić do spłylenia procesów myślowych i nadmiernej zależności od algorytmów, co w konsekwencji obniża zdolność do głębokiego uczenia się i refleksji metapoznawczej. Zagrożeniem jest także nieświadome powielanie błędów i uprzedzeń zawartych w danych, na których trenowane są modele AI. Dlatego niezbędne jest wdrażanie strategii ochronnych, obejmujących m.in. systematyczną autorefleksję, weryfikację źródeł i wyników, krytyczną analizę procesów generowania treści oraz utrzymywanie równowagi między wykorzystaniem narzędzi cyfrowych a tradycyjnymi formami pracy intelektualnej.

2.4. Głębokie uczenie się – charakterystyka i praktyki wspierające rozwój

W literaturze przedmiotu głębokie uczenie się określane jest jako proces wykraczający poza mechaniczne zapamiętywanie treści i polegający na ich głębokim zrozumieniu, integracji z posiadaną już wiedzą oraz świadomym wykorzystaniu w zróżnicowanych kontekstach. (Marton, F., & Säljö, R., 1976) Jego istotą jest aktywne przetwarzanie informacji oraz zdolność do budowania trwałych i powiązanych struktur poznawczych. W procesie tym nowa wiedza nie zastępuje, lecz rozwija istniejące schematy umysłowe, nadając im większą spójność i użyteczność. Jest również procesem kumulatywnym, każda kolejna warstwa poznania osadza się na wcześniejszych doświadczeniach, co umożliwia głębsze i bardziej krytyczne rozumienie badanych zjawisk.

Z perspektywy młodych naukowców utrzymanie zdolności do głębokiego uczenia się oznacza konieczność świadomego unikania strategii powierzchownych, które redukują wysiłek poznawczy i sprzyjają pasywności. Zamiast tego zdolności głębokiego uczenia się można wspierać poprzez kultywowanie autorefleksji, dbanie o ciągłe pogłębianie i łączenie posiadanej wiedzy oraz poszukiwanie nowych kontekstów jej zastosowania. Tak rozumiane uczenie się

staje się nie tylko procesem gromadzenia informacji, lecz także sposobem kształtowania tożsamości badacza. Jest on osobą zdolną do krytycznego, kreatywnego i odpowiedzialnego tworzenia wiedzy w świecie, w którym technologia, w tym generatywna sztuczna inteligencja, stanowi integralny element środowiska poznawczego.

2.5. Iteracja, która łączy

W ujęciu konstruktywistycznym proces uczenia się jest nie tylko aktem przyswajania wiedzy, ale przede wszystkim aktywnym, spiralnym ruchem pomiędzy doświadczeniem, refleksją, tworzeniem i ponownym działaniem. Zarówno spirala kreatywnego uczenia się Mitchela Resnicka, jak i cykl uczenia się oparty na doświadczeniu Davida Kolba odzwierciedlają ten dynamizm. Podkreślają konieczność powtarzania doświadczeń i stopniowego pogłębiania kompetencji. Zarówno u Resnicka, jak i u Kolba kluczowe jest przejście od idei do działania, od działania do refleksji i od refleksji do nowych, udoskonalonych działań. W pracy badawczej ten rytm przybiera postać uniwersalnego cyklu: od formułowania pytań badawczych, przez projektowanie metodologii, zbieranie i analizę danych, po weryfikację, walidację oraz publikację wyników (Creswell 2018).

Generatywna sztuczna inteligencja, odpowiednio włączona do tego procesu, może wzmacniać każdy z jego etapów. W fazie konceptualizacji AI może wspierać generowanie pomysłów, dostarczając różnorodnych perspektyw. Podczas prototypowania i eksperymentowania może pomóc w szybkim testowaniu hipotez, tworzeniu narzędzi czy opracowywaniu założeń. W etapach analizy i walidacji może pełnić rolę czytelnika udzielającego informacji zwrotnej, którą badacz poddaje następnie krytycznej ocenie. Zintegrowanie tych trzech perspektyw: Resnicka, Kolba oraz uniwersalnego cyklu badań ujawnia, że proces uczenia się i proces badawczy są w swej istocie spiralne i iteracyjne. AI nie zastępuje w nich myślenia, lecz może pełnić funkcję katalizatora, przyspieszając przechodzenie pomiędzy kolejnymi etapami i zwiększając głębię refleksji. Według Mitchela Resnicka GenAI nie tyle staje się asystentem tego procesu, ile dodatkowym zasobem, który jest dostępny dla potrzeb pracy badawczej (Resnick, 2025). Ta zmiana perspektywy pozwala uniknąć nadmiernej personalizacji, i jednocześnie wymusza podejście do generowanych treści w podobny sposób, jak do innych wykorzystywanych źródeł.

Tabela 2.1. Integracja modeli Resnicka, Kolba, badań naukowych i roli genAI

Etap Resnicka	Etap Kolba	Etap badań naukowych	Rola Generatywnej AI
Pomysł	Konceptualizacja	Formułowanie pytań badawczych	Proponowanie inspiracji, analiza trendów, generowanie pytań i hipotez

Tworzenie	Konceptualizacja/Doświadczenia	Projektowanie metodologii i narzędzi	Tworzenie prototypów narzędzi, skryptów, modeli. Wspomaganie w opracowaniu procedur badawczych
Zabawa	Doświadczenie /Eksperymentowanie	Zbieranie i analiza danych	Automatyzacja analizy wstępnej; tworzenie kodu; generowanie wstępnych wizualizacji
Dzielenie się wynikami	Refleksja	Weryfikacja i walidacja wyników	Sugestie wyników, wykrywanie niespójności, udzielanie informacji zwrotnej
Refleksja	Refleksja	Publikacja i upowszechnienie wyników	Pomoc w redakcji tekstu, adaptacji językowej, przygotowaniu różnych formatów publikacji
Rekonstrukcja pomysłu	Powrót do konceptualizacji	Generowanie nowych pytań i obszarów badań	Analiza luk w badaniach, generowanie kolejnych kierunków dociekań

Źródło: opracowanie własne

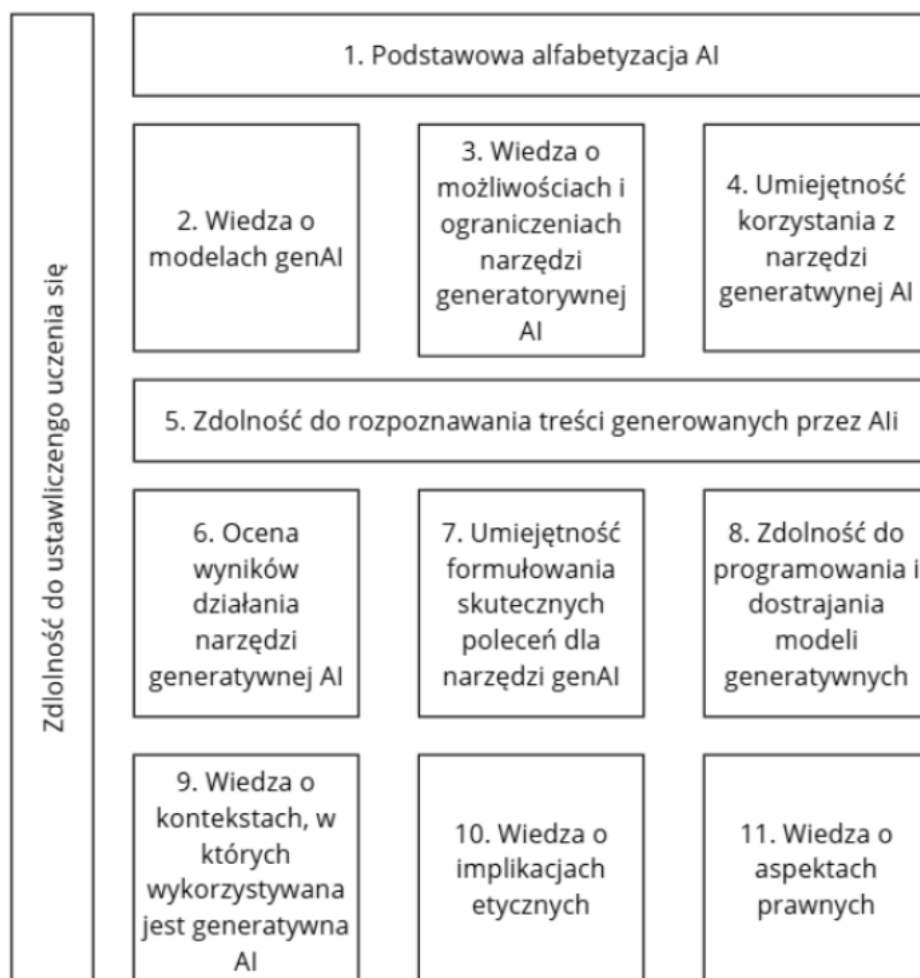
2.6. Kompetencje w zakresie genAI

Kompetencje to funkcja wiedzy, umiejętności i postawy (*European Parliament and the Council, 2006, s. 3*), która jest wynikiem trenowania i rozwijania ludzkich zdolności poznawczych i społecznych. W kontekście rozwoju technologii, w tym generatywnej sztucznej inteligencji, to właśnie kompetencje transwersalne są niezbędne i szczególnie istotne.² To ważny obszar wpływający na zakres i poziom ludzkiego potencjału w posługiwaniu się GenAI. Trudno bez nich realizować zadania naukowe, badawcze i dydaktyczne na wartościowym poziomie.

W celu budowania synergii pomiędzy człowiekiem i technologią konieczne jest rozpoznanie jakie kompetencje z naszej ludzkiej strony są istotne do wzmocnienia, aby ten stan był możliwy do osiągnięcia. Synergia rozumiana jest tutaj jako zjawisko, w którym efekt współdziałania jednostek jest wyższy niż suma efektów osiąganych przez każdą z nich osobno.

Przedstawienie kompetencji warto rozpocząć od tych, które bezpośrednio dotyczą zakresu posługiwania się generatywną AI. Stanowią one odpowiednik alfabetyzacji cyfrowej w zakresie posługiwania się sztuczną inteligencją. W literaturze pojawiła się koncepcja dwunastu kluczowych kompetencji niezbędnych do świadomego wykorzystania GenAI. Jest to zestaw umiejętności, postaw i wiedzy, które musi rozwijać badacz, dydaktyk czy każdy zaawansowany użytkownik AI, aby efektywnie i odpowiedzialnie korzystać z tych narzędzi (Annapureddy i in., 2025). Model zaproponowany przez Annapureddy, Fornaroli i Gatica-Perez (2025) ujmuje kompetencje, począwszy od fundamentalnej orientacji w dziedzinie sztucznej inteligencji, poprzez biegłość w operowaniu narzędziami generatywnymi, aż po świadomą refleksję etyczną nad ich zastosowaniem. Kompetencje te, uporządkowane od podstaw ku coraz bardziej zaawansowanym tworzą logiczny ciąg edukacyjny, przypominający układ spiralny: od

opanowania wiedzy podstawowej, przez rozwój konkretnych umiejętności, po poziom ustawicznego uczenia się, które przenika całość procesu. Taka progresja pokazuje, że skuteczne uczenie się pracy z GenAI odbywa się iteracyjnie, warstwa po warstwie, zgodnie z duchem modeli Kolba i Resnicka, gdzie refleksyjne doświadczenie prowadzi do ciągłego pogłębiania kompetencji.



Rys. 2.3. Diagram dwunastu kompetencji w zakresie generatywnej AI

Źródło: (Annapureddy i in., 2025)

Powyższe składowe są podstawą wiedzy, umiejętności i postaw, które powinny stanowić bazę dla ustawicznego rozwijania i stanowić fundament dla wykorzystywania genAI w wartościowy sposób. Jednak, aby osiągnąć efekt synergii potrzebujemy sięgnąć głębiej, do tego co z punktu widzenia ludzi jest ciągle unikatowym i uniwersalnym zasobem kompetencji.

2.7. Kompetencje transwersalne

Wśród wielu typologii kompetencji jedną z najbardziej znanych jest koncepcja 4C, sformułowana przez Partnership for 21st Century Skills (Waszyngton D.C., 2003). Model ten

obejmuje cztery kluczowe obszary: umiejętność współpracy (*collaboration*), krytyczne myślenie (*critical thinking*), kreatywność (*creativity*) oraz komunikację (*communication*). W literaturze spotkać można również inne propozycje klasyfikacji, które uzupełniają lub modyfikują ten zestaw. Przykładem jest model „Skills for Learning” opracowany przez CEDEFOP, który, obok krytycznego myślenia i kreatywności, wyróżnia także rozwiązywanie problemów (*problem-solving*), adaptacyjność (*adaptability*), cyfrową biegłość (*digital agility*) oraz nastawienie na ustawiczne uczenie się (*lifelong learning mindset*) (CEDEFOP, 2024).

Niezależnie od przyjętej typologii, wspomniane modele kompetencji charakteryzuje szerokie ujęcie, obejmujące różne aspekty ich przejawiania się na poziomie poznawczym, społecznym i praktycznym. Mają one również wspólne cechy, które można określić mianem transgresywności (Nowotny, 2000) i transwersywności (CEDEFOP, 2023). Charakteryzują się przenikaniem zakresów poszczególnych kompetencji, ich wzajemnym wzmacnianiem się oraz wykorzystywaniem wspólnych zasobów w toku realizacji złożonych działań, zarówno w kontekście indywidualnym, jak i zespołowym. W praktyce oznacza to, że kompetencje nie funkcjonują w izolacji, lecz tworzą dynamiczne i złożone systemy, w których zmiana w jednej z nich wpływa na zakres i możliwości pozostałych.

Rozumienia kompetencji transwersalnych nie ogranicza się jedynie do sfery współpracy i relacji międzyludzkich, ale obejmuje także zdolność do efektywnej realizacji zadań we współdziałaniu z cyfrową technologią, w postaci różnych rozwiązań generatywnej sztucznej inteligencji. Praca z GenAI wymaga bowiem nie tylko kompetencji komunikacyjnych, rozumianych tu jako zdolność do precyzyjnego formułowania zapytań i interpretowania odpowiedzi maszyny, lecz także umiejętności kooperacji w ramach realizacji określonych celów, kreatywnego i nieszablonowego rozwijania i realizowania projektów oraz krytycznej ewaluacji otrzymanych rezultatów. Tak rozumiane kompetencje mogą przybierać przedstawione poniżej wymiary:

- komunikacja jest interakcją człowieka z systemem opartym na modelach językowych;
- współpraca jest procesem wspólnego konstruowania wiedzy w oparciu o propozycje generowane przez algorytmy;
- kreatywność jest zdolnością do twórczego wykorzystania treści wytworzonych przez AI jako materiału do dalszej pracy;
- krytyczne myślenie jest koniecznością nieustannej weryfikacji poprawności, spójności i wiarygodności danych dostarczanych przez technologię.

W tym kontekście transwersywność kompetencji obejmuje zarówno przekraczanie granic dyscyplin i form myślenia, jak i poszerzanie przestrzeni działania poprzez synergiczne łączenie ludzkiego potencjału z możliwościami maszyn. Taka perspektywa może powodować, że skracając czas i wysiłek w pewnych powtarzalnych i nużących zakresach pracy, badacze powinni zwiększać zaangażowanie w aspekty, które pozwolą przenosić nasze projekty na wyższy poziom, również dzięki nieszablonowemu podejściu do genAI.

Ponieważ kompetencje transwersalne są dynamicznym i powiązanim ze sobą systemem, na który wpływ mają nasze świadome i nieświadome doświadczenia, ważne jest takie kierowaniem procesem uczenia się, aby mogły być one utrwalane i rozszerzane. Włączanie genAI do rytuałów i procesów pracy badawczej pozostaje w sprzężeniu z wiedzą, umiejętnościami i postawami, które w ustawicznych procesach rozwoju poddajemy modyfikacji. Stąd potrzeba refleksyjnego monitorowania wpływu korzystania z tych narzędzi na warsztat młodych dydaktyków i badaczy.

2.8. Model SAMR

Model SAMR³, zaproponowany przez Rubena Puentedurę (2010), to czteropoziomowa struktura opisująca sposób korzystania z technologii w procesie edukacyjnym: *substitution* (zastąpienie), *augmentation* (rozszerzenie), *modification* (modyfikacja) oraz *redefinition* (redefinicja). W kontekście pracy doktorantów i młodych badaczy z generatywną sztuczną inteligencją, model ten może pełnić funkcję narzędzia służącego refleksji i rozpoznaniu poziomu korzystania z GenAI w pracy dydaktycznej i badawczej. Pozwala uchwycić, na jakim etapie integracji technologii jesteśmy oraz jak rozwój naszych kompetencji jest sprzężony ze stosowaniem narzędzi AI na wyższych poziomach zaawansowania.

W perspektywie kompetencji transwersalnych, każdy z poziomów modelu SAMR charakteryzuje się głębszym zakresem ich wykorzystania. Na najniższym poziomie zastąpienia, rola GenAI sprowadza się do pełnienia funkcji cyfrowej wersji tradycyjnych narzędzi, np. automatycznego zapisu pomysłów, tłumaczenia materiałów czy wyszukiwania informacji. Krytyczne myślenie jest tu stosowane w wąskim zakresie, ograniczając się do podstawowej oceny poprawności treści; kreatywność ma charakter operacyjny; komunikacja jest prosta, a możliwość współpracy incydentalna.

Przy poziomie rozszerzenia, zastosowanie technologii może wspierać pogłębioną analizę i optymalizację działań badawczych. Krytyczne myślenie obejmuje porównywanie źródeł i ich weryfikację; kreatywność polega na generowaniu alternatywnych rozwiązań; komunikacja staje

się bardziej intencjonalna, a współpraca świadomym procesem wymiany danych i interpretacji wyników.

Na poziomie modyfikacji, badacz sięga po GenAI, aby integrować ją z procesami badawczymi i twórczymi w sposób strukturalny. Jej wykorzystanie wspiera modelowanie problemów, symulacje oraz generowanie dostępnych scenariuszy badawczych. Krytyczne myślenie obejmuje analizę konsekwencji i wybór optymalnych strategii, kreatywność przejawia się w tworzeniu nowych koncepcji badawczych, komunikacja nabiera charakteru argumentacyjnego i poddającego w wątpliwość, a współpraca wymiaru koordynacji zadań i rozwiązywania złożonych problemów.

Na najwyższym poziomie redefinicji, możliwe staje się projektowanie działań badawczych i dydaktycznych o innowacyjnym charakterze, poszerzających pola badawcze, które bez dostępu do rozwiązań technologicznych byłyby niedostępne i niemożliwe do wdrożenia. Kompetencje 4C osiągają tu najszerszy zakres: krytyczne myślenie obejmuje metapoznanie oraz ewaluację wpływu np. na otoczenie społeczne, kreatywność przyjmuje formę transformacyjną, komunikacja ma charakter wielokanałowy i międzykulturowy uwzględniający różne perspektywy, a współpraca opiera się na otwartości i ustawicznej wymianie. Synteza powyższego opisu znajduje się w tabeli 2.2.

Tabela 2.2. Poziomy kompetencji w modelu SAMR

Poziom SAMR	Krytyczne myślenie	Kreatywność	Komunikacja	Współpraca
Zastąpienie	Podstawowa ocena poprawności i aktualności treści	Operacyjne dostosowanie istniejących pomysłów	Proste przekazywanie informacji	Opcjonalna, doraźna współpraca
Rozszerzenie	Porównywanie źródeł, weryfikacja danych	Generowanie alternatywnych rozwiązań, optymalizacja	Świadoma wymiana danych i opinii	Współdzielenie zadań i wstępna koordynacja
Modyfikacja	Analiza konsekwencji, wybór strategii, modelowanie problemów	Projektowanie nowych koncepcji badawczych	Argumentacyjna, wieloetapowa komunikacja	Koordynacja i wspólne rozwiązywanie problemów
Redefinicja	Metapoznanie, refleksja etyczna i społeczna	Kreatywność transformacyjna, innowacje przełomowe	Wielokanałowa, międzykulturowa komunikacja	Otwartość i elastyczna wymiana

Źródło: opracowanie własne

Rozwój własnych kompetencji badawczych i gromadzenie doświadczeń pracy zespołowej z ludźmi i maszynami, są warunkiem przechodzenia na wyższy poziom modelu SAMR w pracy z GenAI. To właśnie progresja w obrębie 4C pozwala zbliżyć się do modelu synergii człowiek - maszyna.

2.9. Etyczny wymiar synergii

W kontekście rozwijania warsztatu pracy dydaktycznej i badawczej ważne jest, abyśmy nie pominęli aspektu etycznego. Będzie on pojawiał się też w innych miejscach tej publikacji, dlatego w tym miejscu wskazane zostaną dwie ogólne perspektywy, które mogą dać uniwersalny i szeroki fundament dla tworzenia warunków dla synergii.

2.10. AI Act

Akt o sztucznej inteligencji Unii Europejskiej (AI Act, 2024) stanowi pierwszy w skali globalnej kompleksowy dokument prawny regulujący obszar sztucznej inteligencji. Pomimo, że akt jest zawiłym i obszernym dokumentem, zostało z niego wyłonione to, co istotne i uniwersalne dla naszych rozważań. Jego fundamentem jest ochrona wartości podstawowych, które mają być gwarantem, że rozwój i wdrażanie technologii AI nie będą sprzeczne z zasadami demokratycznego ładu i prawami człowieka. Dokument odwołuje się do katalogu praw podstawowych Unii Europejskiej, obejmujących między innymi prawo do godności ludzkiej, wolności i równości, ochrony danych osobowych i prywatności, zakazu dyskryminacji, prawa do sprawiedliwego procesu, wolności wypowiedzi, a także prawa do edukacji i uczestnictwa w życiu społecznym i kulturalnym. Regulacja ta podkreśla, że celem nie jest zahamowanie innowacyjności, lecz zapewnienie, aby postęp technologiczny wzmacniał fundamenty życia społecznego, a nie je podważał. W tym sensie AI Act jest dokumentem o charakterze nie tylko prawnym, lecz także głęboko aksjologicznym, wskazującym na konieczność rozwijania technologii w zgodzie z wartościami, które konstytuują europejską tożsamość. Ta perspektywa pozwala pokazać, że w synergii pomimo dodawania potencjału człowieka i technologii, to właśnie perspektywa człowieka jest tym, o co należy dbać i co należy rozwijać.

2.11. Linie na piasku

Ten humanocentryczny punkt widzenia i wytyczne do jego wdrażania w bardziej przystępny niż Acta AI sposób opracował zespół Mitchela Resnicka (Resnick, 2025). Powstały one dla potrzeb rozwijania ekosystemów edukacyjnych opracowanych przez grupę Lifelong Kindergarten i Fundację Scratcha. Są to wytyczne dotyczące wartości i ograniczeń nazwane poetycko: *Guiding Stars* i *Lines in the sand*, a przedstawiają się w następujący sposób (Resnick, 2025)⁴:

Gwiazdy przewodnie (*Guiding Stars*):

- kreatywność: wspieranie w wyrażaniu siebie;
- sprawczość: utrzymywanie w rękach człowieka mocy decydowania i dokonywania wyborów;
- równość: zapewnienie bezpłatnego, sprawiedliwego, inkluzywnego i dostępnego charakteru;

– wspólnota: podkreślanie znaczenia ludzkich więzi i współpracy.

Linie na piasku (*Lines in the sand*):

– bezpieczeństwo: zapobieganie szkodliwym treściom;

– etyczność: ochrona danych użytkowników;

– transparentność: dostęp do informacji o tym, jak działa technologia;

– humanocentryczność: brak podważania ludzkiej sprawczości ani relacji międzyludzkich.

Wymienione punkty wytyczają, podobnie jak Acta AI ramy, które wyznaczają ramy dla twórców technologii opartej na AI. Jednocześnie wskazują też jakich rozwiązań należy szukać i na co zwracać uwagę przy realizacji projektów, które włączają generatywną sztuczną inteligencję w pracy dydaktycznej i naukowej.

Podsumowanie rozdziału 2

Przedstawione w tym rozdziale analizy ukazują uczenie się jako zjawisko oraz dynamiczny i wielowymiarowy proces. Odwołanie do modeli Kolba i Resnicka pozwoliło dostrzec, iż nie ogranicza się on do linearnego przyswajania wiedzy, lecz opiera się na cykliczności, refleksji oraz przetwarzaniu doświadczeń. Proces uczenia się znajduje w tym sensie swój odpowiednik w praktykach badawczych, które również rozwijają się w rytmie powtarzającego się cyklu pytań, metod i weryfikacji.

Wprowadzenie perspektywy kompetencyjnej ukazało, że rozwój w obszarze generatywnej sztucznej inteligencji nie może być sprowadzony do nabywania biegłości technicznej, ale wymaga szerszej refleksji nad kompetencjami przekrojowymi. Model kompetencji transwersalnych wskazuje, że kluczowe dla efektywnego i odpowiedzialnego wykorzystywania nowych technologii są: krytyczne myślenie, kreatywność, komunikacja oraz współpraca. Kompetencje te nie występują w izolacji, lecz pozostają ze sobą w relacjach wzajemnego przenikania i wzmocnienia, co potwierdza ich systemowy charakter.

Model SAMR pozwolił pokazać, że wykorzystanie generatywnej AI w badaniach i edukacji ma charakter progresywny i poziom ich zastosowania zależy od rozwoju kompetencji badaczy. Na poziomach podstawowych narzędzie pełni funkcję wspierającą i uzupełniającą, natomiast wraz z rozwojem kompetencji otwiera możliwość redefiniowania praktyk badawczych i dydaktycznych. W tym kontekście AI staje się zasobem, którego wartość zależy od zdolności badacza do krytycznej refleksji, twórczego eksperymentowania i świadomego zarządzania procesami poznawczymi.

Rozwój ten może prowadzić do budowania synergii człowieka i maszyny. Synergia ta polega na wzajemnym uzupełnianiu się ludzkich kompetencji poznawczych, społecznych i etycznych oraz potencjału obliczeniowego i generatywnego AI. Proces uczenia się i rozwój kompetencji stają się tu warunkiem umożliwiającym wartościowe współdziałanie, które nie redukuje człowieka do roli pasywnego użytkownika, lecz czyni go aktywnym współtwórcą. Kwestie etyczne są nieodłącznym elementem tej synergii, przypominając, że budowanie warsztatu badawczego opartego na AI wymaga nie tylko sprawności technicznej i poznawczej, ale także odpowiedzialności, świadomości konsekwencji oraz troski o dobro wspólne. Tak rozumiane sprzężenie wiedzy, kompetencji i technologii otwiera młodemu badaczowi drogę do twórczego i wartościowego praktykowania nauki, w której generatywna sztuczna inteligencja staje się czynnikiem sprawczym w procesie pogłębiania procesu myślenia i rozumienia, a nie jego substytutem.

3. Ograniczenia i zagrożenia wykorzystywania AI w pracy naukowej

Dynamiczny rozwój sztucznej inteligencji w nauce przynosi nie tylko korzyści, lecz także szereg wyzwań, które wymagają świadomego i krytycznego podejścia. Korzystanie z modeli generatywnych wiąże się z ryzykiem halucynacji, błędnych cytowań, stronniczości danych oraz problemów etycznych i prawnych. W niniejszym rozdziale omówiono najważniejsze zagrożenia towarzyszące integracji AI z procesem badawczym, wskazując zarówno ich źródła, jak i potencjalne konsekwencje dla rzetelności nauki oraz wiarygodności publikacji.

3.1. Halucynacje i stronniczość

Generatywna AI zwłaszcza duże modele językowe coraz częściej wspiera pisanie tekstów, przeglądy literatury, analizę danych i projektowanie eksperymentów. Zyski produktywności są realne (Khalifa, Albadawy, 2024), ale ryzyka dla rzetelności naukowej wymagają ścisłej kontroli, w szczególności halucynacji, czyli zmyślenia faktów i stronniczości (bias) (Mittelstadt, Russell, Wachter, 2023; Ferrara, 2023).

LLM potrafią generować płynny, poprawny stylistycznie tekst, który bywa merytorycznie błędny. Udokumentowano zarówno wymyślanie faktów, jak i fikcyjne cytowania nieistniejące prace lub zniekształcone dane bibliograficzne (Alkaissi, McFarlane, 2023; Walters, Wilder, 2023).

Halucynacje są szczególnie niebezpieczne w streszczeniach i miniprzeglądach, gdzie model może nadawać badaniom zbyt szerokie wnioski lub pewność twierdzeń. Dlatego LLM powinny służyć najwyżej jako asystenci językoworedakcyjni między innymi w celu parafrazowania, tłumaczenia, a nie jako autonomiczne źródła wiedzy (Mittelstadt i in., 2023). W związku tym, należy zawsze weryfikować fakty i sprawdzać każdą referencję w katalogach, bazach - DOI, czasopismo, rocznik.

W pracach z udziałem AI należy ujawnić zakres jej użycia oraz zachować krytyczną ocenę treści (Alkaissi, McFarlane, 2023; Walters, Wilder, 2023).

Bias (stronniczość) może wynikać z danych treningowych, na przykład nadreprezentacji języka angielskiego, określonych regionów, architektury i polityk modelu. Skutkiem są deformacje, faworyzowanie określonych nurtów, pomijanie mniejszościowych perspektyw i utrwalanie stereotypów (Ferrara, 2023). W pisaniu i przeglądach literatury może to oznaczać anglocentryczność zestawień, przewagę pism wysokoimpaktowych oraz nadmierne potwierdzanie konsensusu (Ferrara, 2023).

W analizie danych i wspomaganiu eksperymentów bias prowadzi do systematycznych błędów modeli oraz decyzji projektowych, które konserwują dominujące paradygmaty (Ferrara, 2023). Warto zatem dbać o różnorodne źródła (języki, regiony, szkoły badawcze), testować modele pod kątem sprawiedliwości i wprowadzać nadzór człowieka „human-in-the-loop” na kluczowych etapach (Ferrara, 2023; Mittelstadt i in., 2023).

Modele „czarnej skrzynki” utrudniają wyjaśnialność wyników i mogą nadmiernie dopasowywać się do wzorców w danych, co osłabia uogólnialność. LLM i inne modele generatywne mogą sugerować technicznie nieadekwatne lub z etycznego punktu widzenia wątpliwe kroki eksperymentalne, jeśli brak im kontekstu z danej tematyki (Mittelstadt i in., 2023).

Warto więc łączyć metody AI z procedurami weryfikacji na przykład z walidacją na danych zewnętrznych, analizą wrażliwości, dokumentowaniem danych i decyzji, a propozycje AI traktować jako hipotezy do sprawdzenia, nie rozstrzygnięcia (Khalifa, Albadawy, 2024; Mittelstadt i in., 2023).

Kluczowe ramy działania związane halucynacjami i stronniczością:

1. Rola AI: warto używać LLM głównie do stylu, parafrazy i tłumaczenia, nie do tworzenia tez, wniosków (Mittelstadt i in., 2023).
2. Weryfikacja: należy sprawdzać każde twierdzenie i cytowanie; nie kopiować bibliografii z AI bez kontroli (Walters, Wilder, 2023; Alkaissi, McFarlane, 2023).

3. Przejrzystość: powinno się ujawnić użycie AI w metodyce, podziękowaniach; autor bierze pełną odpowiedzialność (Khalifa, Albadawy, 2024).
4. Sprawiedliwość i uczciwość: warto zapewnić różnorodność źródeł i testować modele pod kątem biasu (Ferrara, 2023).
5. Nadzór i replikowalność: powinno się dokumentować dane, parametry i wersje narzędzi; wyniki AI replikować niezależnymi metodami (Mittelstadt i in., 2023).

AI może podnosić produktywność badań, ale bez rygorów weryfikacji i kontroli biasu zagraża rzetelności nauki. Krytyczny, przejrzysty i odpowiedzialny użytek z człowiekiem w roli decydenta pozostaje warunkiem bezpiecznej integracji AI w warsztacie badawczym (Mittelstadt i in., 2023; Ferrara, 2023).

3.2. Skala błędnych cytowań generowanych przez AI

Upowszechnienie dużych modeli językowych (LLM) w środowisku akademickim przyniosło szybkie wsparcie dla przeglądów literatury i pisanie, ale jednocześnie ujawniło poważny problem halucynacji bibliograficznych: generowania pozornie wiarygodnych, a w rzeczywistości fałszywych lub zniekształconych pozycji w spisie literatury. Mechanizm działania LLM polega na przewidywaniu kolejnych słów, a nie na weryfikacji faktów w bazach źródłowych, dlatego błędy w danych identyfikacyjnych (np. DOI, PMID, tom, strony) są częste i systematyczne (Alkaissi, McFarlane, 2023). Zjawisko to uderza w rdzeń rzetelności naukowej, ponieważ uniemożliwia replikację i kontrolę źródeł, a przez to osłabia zaufanie do całych wywodów naukowych (OrduñaMalea, CabezasClavijo, 2023; Walters, Wilder, 2023).

Skala problemu została dobrze udokumentowana w badaniach z ostatnich lat. W medycynie wykazano, że na 115 pozycji bibliograficznych wygenerowanych przez LLM 47% stanowiły cytowania całkowicie zmyślane, kolejne 46% zawierało istotne błędy, a tylko 7% było w pełni poprawnych (Bhattacharyya i in., 2023). W innym badaniu odsetek sfabrykowanych odwołań przekraczał połowę wszystkich pozycji dołączanych do odpowiedzi klinicznych, mimo że same odpowiedzi bywały częściowo merytorycznie trafne (Gravel, D'AmoursGravel, Osmanliu, 2023).

Analizy porównawcze sugerują również, że nowsze modele ograniczają skalę halucynacji, lecz jej nie eliminują, dla GPT3.5 odsetek fałszywych cytowań był większościowy, podczas gdy dla GPT4 spadał do poziomu rzędu kilkunastu procent (Day, 2023; Walters, Wilder, 2023). W obszarach niszowych (np. wybrane tematy geograficzne) raportowano szczególnie wysokie natężenie „widmowych” referencji (Day, 2023), a w praktyce klinicznej radiologii jedynie około jedna trzecia wygenerowanych pozycji dawała się

zweryfikować i adekwatnie wspierała odpowiedź (Wagner, ErtlWagner, 2024).

Dodatkowe obserwacje wskazują na częste częściowe halucynacje, istnienie publikacji przy jednoczesnych błędach w numerach DOI/PMID czy metadanych (Athaluri i in., 2023).

W literaturze bibliometrycznej opisano przy tym narastające zjawisko „ghost references”, referencji widmowych, które wyglądają na rzetelne (prawidłowo sformatowane nazwiska, tytuły, nazwy czasopism), lecz nie istnieją w żadnej bazie (Orduña Malea, Cabezas Clavijo, 2023). Ryzyko ich przenikania do obiegu naukowego rośnie zwłaszcza w preprintach oraz w publikacjach o słabszych procedurach redakcyjnych. W odpowiedzi na te wyzwania zaproponowano miary oceny zjawiska na przykład Reference Hallucination Score (RHS), aby obiektywnie porównywać narzędzia i śledzić postęp redukcji błędów (Aljamaan i in., 2024).

Konsekwencje dla praktyki badawczej są wielowymiarowe. Fałszywe lub zniekształcone cytowania wprowadzają dezinformację do przeglądów literatury, utrudniają śledzenie źródeł, a także mogą stać się mimowolnym „sygnałem ostrzegawczym” dla recenzentów, wskazującym na udział generatywnej AI w tworzeniu tekstu (Bhattacharyya i in., 2023; Walters, Wilder, 2023). W skrajnych przypadkach prowadzą do korekt lub wycofań manuskryptów, co generuje koszty reputacyjne i organizacyjne.

Zalecenia operacyjne (protokół minimalny) dla autorów i redakcji:

- należy weryfikować istnienie każdej pozycji w niezależnych bazach (Crossref, PubMed, Scopus/Web of Science, katalogi bibliotek);
- należy krzyżowo sprawdzać metadane (autorzy, rok, tytuł, czasopismo, tom, strony, DOI/PMID) i zgodność treści artykułu z przytaczanym twierdzeniem;
- należy unikać proszenia LLM o gotowe bibliografie; dopuszczalne jest traktowanie sugestii jako punktu wyjścia, wyłącznie po pełnej weryfikacji;
- należy dokumentować proces walidacji źródeł (na przykład w materiałach uzupełniających) oraz stosować wskaźniki jakości, takie jak RHS (Reference Hallucination Score, to wskaźnik oceniający autentyczność i relewancję cytatów generowanych przez modele językowe w celu pomiaru ryzyka halucynacji), do oceny ryzyka halucynacji i porównywania narzędzi (Aljamaan i in., 2024);
- należy wdrażać w redakcjach procedury kontroli referencji (automatyczne sprawdzanie DOI/PMID, wyłapywanie nieistniejących tytułów), a w pracach dyplomowych i doktorskich - formalny wymóg potwierdzenia źródeł.

Wykorzystanie LLM w pracach naukowych wymaga podejścia defensywnego: narzędzia mogą przyspieszać pracę koncepcyjną i językową, jednak aparatu bibliograficznego

nie wolno delegować na modele generatywne. Dopiero rygorystyczna weryfikacja każdej referencji przywraca minimalny poziom zaufania do tekstu (Alkaissi, McFarlane, 2023; Orduña Malea, Cabezas Clavijo, 2023; Walters, Wilder, 2023).

3.3. Zagadnienia dotyczące etyki

Dynamiczny wzrost użycia narzędzi AI, w tym generatywnych modeli językowych w badaniach przyspiesza analizy, lecz rodzi ryzyka dla integralności naukowej: niejasne autorstwo, ograniczona przejrzystość modeli, uprzedzenia algorytmiczne, rozmyta odpowiedzialność i zagrożenia dla prywatności (Resnik, Hosseini, 2024).

Konsensus redakcyjny i etyczny jest spójny: AI nie spełnia kryteriów autorstwa i nie powinna być wskazywana jako współautor, bo nie może odpowiadać za treść, deklarować konfliktów interesów ani zatwierdzać ostatecznej wersji pracy (Flanagin i in., 2023; Resnik Hosseini, 2024). Dopuszczalne jest korzystanie z AI jako narzędzia, o ile jawnie ujawnimy zakres i cel użycia (nazwa modelu, sposób wsparcia, istotne ustawienia) coraz więcej czasopism wprost tego wymaga (Yoo, 2025). Autorzy pozostają odpowiedzialni za treść, w tym za zapobieganie niezamierzonemu plagiatowi czy „ghostwritingowi” AI (Flanagin i in., 2023; Resnik, Hosseini, 2024).

Zaawansowane modele z problemem tak zwanej „czarnej skrzynki” utrudniają reprodukowalność i audytowalność wyników. Etyka badań wymaga zatem maksymalnej przejrzystości: ujawniania danych, protokołów, metryk oraz ograniczeń narzędzia, a tam, gdzie to możliwe stosowania metod XAI (explainable AI) i dokumentowania decyzji modelu (BouhouitaGuermech i in., 2023; Resnik, Hosseini, 2024). Transparentność powinna obejmować nie tylko to, jak model działa, ale i jaką rolę odgrywa w toku badania (BouhouitaGuermech i in., 2023).

Modele uczą się z danych, więc dziedziczą i wzmacniają istniejące biasy (stronniczości), co może zniekształcać wnioski i utrzymywać nierówności (BouhouitaGuermech i in., 2023; Ateriya i in., 2025). Minimalizacja ryzyka wymaga: walidacji na zróżnicowanych podzbiorach, monitorowania metryk sprawiedliwości, stosowania technik debiasing oraz krytycznego przeglądu wyników przez zespół o zróżnicowanych kompetencjach (Ateriya i in., 2025; Resnik, Hosseini, 2024). Etyka obejmuje też sprawiedliwy dostęp do narzędzi AI, aby nie pogłębiać luki między instytucjami (Ateriya i in., 2025).

Odpowiedzialność zawsze spoczywa na ludziach, AI nie jest podmiotem moralnym ani prawnym (Resnik, Hosseini, 2024). W praktyce oznacza to stały nadzór człowieka, krytyczną weryfikację wyników AI i gotowość do ich uzasadnienia wobec recenzentów i opinii publicznej.

Redakcyjne wytyczne podkreślają, że autorzy muszą sprawdzać poprawność i rzetelność treści wygenerowanych przez modele, bo te mogą brzmieć autorytatywnie, lecz być błędne lub stroniczne (Yoo, 2025; Resnik & Hosseini, 2024).

Korzystanie z usług AI zwłaszcza w chmurze wiąże się z ryzykiem wycieku treści nieopublikowanych, danych osobowych i informacji poufnych. Etyka wymaga minimalizacji danych, anonimizacji, starannego doboru platform, gdzie preferencyjnie są rozwiązania lokalne, instytucjonalne, a także jasnego informowania uczestników o użyciu AI w analizie ich danych (BouhouitaGuermech i in., 2023; Resnik, Hosseini, 2024).

3.4. Prawa autorskie RODO/GDPR, odpowiedzialność autorów

W ostatnich latach generatywna AI (między innymi ChatGPT) szybko weszła do warsztatu badawczego, ale redakcje i środowiska etyczne zgodnie wymagają ujawniania użycia narzędzi i podkreślają, że odpowiedzialność za treść spoczywa wyłącznie na autorach ludziach (Ganjavi i in., 2024; Hosseini, Resnik, Holmes, 2023). Najnowsze przeglądy polityk czasopism wskazują jednocześnie na duże zróżnicowanie praktyk oraz ograniczoną przydatność detektorów AI do decyzji redakcyjnych stąd nacisk na jawność i weryfikację merytoryczną przez autorów (Yoo, 2025). W literaturze toczy się też debata, na ile odpowiedzialność powinna być warunkiem autorstwa, ale konsensus redakcyjny pozostaje: AI nie spełnia kryteriów autorstwa (Levy, 2025).

Badania nad LLM pokazują, że modele potrafią odtwarzać fragmenty danych treningowych, w tym utwory chronione w odpowiedzi na zapytania, co stwarza realne ryzyko naruszeń praw autorskich w publikacjach naukowych (Carlini i in., 2021). Nowsze prace porównawcze sugerują, że poziom „copyright compliance” różni się istotnie między modelami i scenariuszami użycia, a mechanizmy obronne bywają niespecyficzne (Mueller i in., 2024). Dla autorów oznacza to konieczność kontroli wyjść modelu (na przykład porównanie ze źródłami, własny plagiatcheck wygenerowanych fragmentów) i dokumentowania procesu (prompty, wersje narzędzi, zakres edycji). (Carlini i in., 2021; Mueller i in., 2024).

Analizy prawne podkreślają, że szkolenie modeli na utworach chronionych budzi wątpliwości w świetle wyjątków TDM (*Text Data Mining*) oraz obowiązków przejrzystości i licencjonowania; zarówno prace naukowe, jak i opracowania prawnicze wskazują na luki i rozbieżności interpretacyjne (Novelli i in., 2024; Li, 2024; Quintais, 2025). Część autorów argumentuje, że szkolenie genAI nie mieści się w wyjątkach TDM i co do zasady wymaga licencji (Dornis, 2024). W praktyce badawczej bezpieczną strategią jest wykorzystywanie zbiorów z jasnymi licencjami, unikanie wklejania do promptów treści cudzych chronionych

prawem oraz dodanie noty o weryfikacji i oryginalności (z określeniem narzędzi i wersji). (Novelli i in., 2024; Li, 2024; Quintais, 2025; Dornis, 2024).

W kontekście LLM przetwarzanie masowych, publicznie dostępnych danych osobowych (na przykład postów) do trenowania, wdrażania modeli wymaga ważnej podstawy prawnej; sama publiczna dostępność nie znosi obowiązków z art. 6 RODO (Kuru, 2024). Zwraca się uwagę na sporne stosowanie uzasadnionego interesu jako podstawy oraz na konsekwencje dla późniejszego użycia modelu, jeśli etap rozwoju był niezgodny z prawem (Ruscheimer, 2025). Dodatkowym wyzwaniem są prawa osób (dostęp, sprostowanie, usunięcie) wobec systemów trudnych do oduczenia, w literaturze proponuje się kombinację środków organizacyjnych i technicznych: logowanie danych wejściowych, ograniczanie wprowadzania danych osobowych do promptów (Ruscheimer, 2025; Kuru, 2024).

Przeglądy wskazują na użyteczność technik prywatności: różnicowa prywatność, uczenie federacyjne, SMC (Secure Multi-Party Computation, czyli bezpieczne obliczenia wielostronne), homomorficzne szyfrowanie w kontekście generatywnej AI, choć z wyraźnymi kompromisami jakościowymi i kosztowymi (Feretzakis i in., 2024). Dla zespołów badawczych praktyczne minimum to: DPIA (*Data Protection Impact Assessment*) dla projektów z danymi osobowymi, ograniczanie danych w promptach, kontrola dostępu do logów i modeli, a także wyraźne informacje o ochronie prywatności wobec uczestników badań. (Feretzakis i in., 2024).

Przeglądy polityk najwyżej notowanych czasopism pokazują, że niemal wszystkie zakazują podawania AI jako autora i wymagają jawności użycia narzędzi, nierzadko z wyspecyfikowaniem: nazwy, wersji, zakresu, lokalizacji ujawnienia (metody, podziękowania, cover letter) (Ganjavi i in., 2024; Yoo, 2025). Etyczne opracowania akcentują obowiązek autorów odnośnie krytycznej oceny wkładu AI, opisanie ograniczeń i ryzyk oraz zapewnienia braku plagiatu (Hosseini i in., 2023; Resnik, 2024). Jednocześnie badania empiryczne dokumentują halucynacje w tym fałszywe cytowania co czyni ludzką weryfikację bibliografii obowiązkową (Chelli i in., 2024; Farquhar i in., 2024).

Podsumowanie rozdziału 3

Przedstawione w tym rozdziale aspekty, wskazują że wykorzystanie sztucznej inteligencji w pracy naukowej wymaga zachowania ostrożności i stosowania ścisłych procedur kontrolnych. AI może wspierać badaczy, ale generuje również poważne ryzyka od halucynacji treści i fałszywych cytowań, po naruszenia etyczne, prawne i związane z ochroną danych. Kluczem do bezpiecznej integracji tych narzędzi pozostaje krytyczna weryfikacja wyników, przejrzystość w ujawnianiu ich użycia oraz utrzymanie pełnej odpowiedzialności po stronie

badacza. Tylko w takich warunkach możliwe jest wykorzystanie potencjału AI bez osłabiania fundamentów rzetelności naukowej.

4. Zmiany metodyczne z użyciem AI i przyszłe kierunki

Rozwój sztucznej inteligencji w coraz większym stopniu wpływa na metodologię badań naukowych, zmieniając sposoby wyszukiwania, analizy i prezentacji wiedzy. AI wspiera nie tylko pojedyncze zadania, lecz także kształtuje nowe podejścia do przeglądów literatury, projektowania badań oraz procesów recenzyjnych. W tej części przedstawiono, jak generatywna sztuczna inteligencja przekształca tradycyjne praktyki badawcze oraz jakie kierunki zmian można przewidywać w najbliższych latach.

4.1. Zmiana metodologii przeglądów literatury z wykorzystaniem AI

Generatywna AI nie tylko przyspiesza pracę badacza, zmienia wzorce wytwarzania wiedzy od sposobów wyszukiwania i syntezy piśmiennictwa po redagowanie artykułów i praktyki recenzyjne. Zmiana polega na przejściu od pojedynczych, epizodycznych zadań do zintegrowanych, dynamicznych procesów, w których człowiek prowadzi, a systemy AI wspierają planowanie, kontrolę jakości i uaktualnianie treści. Autorzy podręczników i przeglądów poświęconych temu zwrotowi podkreślają, że skuteczne i etyczne wdrożenia wymagają przejrzystości, krytycznej oceny wyników i zachowania odpowiedzialności autorów za całość pracy (Haber, Jemielniak, Kurasiński, Przegalińska, 2025; Cheng, Calhoun i Reedy, 2025).

Przegląd literatury staje się procesem ciągłym i półautomatyzowanym. Redakcje i badacze wskazują na potrzebę dynamicznych przeglądów, aktualizowanych w miarę napływu danych i wspieranych przez narzędzia AI do ekstrakcji, grupowania i streszczania treści, przy jednoznacznym pozostawieniu rozstrzygnięć merytorycznych po stronie ludzi (Tomczyk, Brüggemann i Vrontis, 2024).

Dowody porównawcze pokazują zarówno potencjał, jak i ograniczenia LLM w przeglądach naukowych. W badaniach zestawiających wyniki ChatGPT z Google Scholar uzyskano szerokie pokrycie tematu, ale także niestabilność faktów i ryzyko halucynacji, co przesądza o konieczności weryfikacji źródeł i ścisłej kontroli eksperckiej (Tomczyk, Brüggemann, Mergner i Petrescu, 2024).

W praktyce nowa metodologia łączy trzy warstwy:

1. wspomagane przez AI wyszukiwanie i klasteryzację publikacji,

2. ręczną ocenę jakości i trafności,
3. rejestrowalną ścieżkę decyzji i aktualizacji.

Podręczniki naukowe dotyczące pracy z AI w nauce zalecają między innymi iteracyjne sprawdzanie treści i cytowań, porównywanie streszczeń wielu modeli, a przede wszystkim porządkowanie i weryfikację bibliografii przed dodaniem do publikacji (Haber, Jemielniak, Kurasiński i Przegalińska, 2025).

Równolegle rośnie znaczenie list kontrolnych i wytycznych etycznych. Zespół badawczy Cheng, Calhoun i Reedy (2025) proponuje praktyczny zestaw pytań kontrolnych dla autorów dotyczących wkładu intelektualnego, kompetencji badawczych, rzetelności treści i jawności użycia narzędzi, co łatwo przełożyć na etap planowania i raportowania przeglądów (Cheng, Calhoun i Reedy, 2025).

Analizy metodologiczne uzupełnia kartografia nauki: mapowanie zmiennych i relacji między polami badawczymi, które wsparte algorytmami wzmacniają przejrzystość ram pojęciowych i decyzji inkluzyjnych w przeglądzie (Tomczyk, Brüggemann i Paul, 2024).

W literaturze popularnonaukowej i zarządczej akcent kładzie się na współtworzenie i demokratyzację, otwarte repozytoria, współpracę interdyscyplinarną oraz łączenie ludzkiej wnikliwości z przetwarzaniem masowym, co sprzyja bardziej inkluzywnym i aktualnym syntezom (Przegalińska i Triantoro, 2021; Przegalińska i Jemielniak, 2023).

Wreszcie, z innej perspektywy szerokie omówienia roli AI w akademii podkreślają, że transformacja przeglądów literatury jest częścią głębszej zmiany epistemicznej - narzędzia stają się współautorami procesów badawczych, ale odpowiedzialność metodologiczna i interpretacyjna pozostaje po stronie naukowców (Beshr, Ateeq, Ateeq i Alaghbari, 2025).

4.2 Przyszłe kierunki zmian związane z pisanem publikacji i recenzowaniem

W pisaniu i publikowaniu rozwija się model „asystenta redakcyjnego”, czyli system, który pomaga w strukturze, dopracowaniu stylu, dopasowaniu do wymogów czasopisma, a nawet w symulacjach perspektywy recenzenta; kontrola jakości, oryginalność i decyzje autorstwa pozostają jednak ludzkie (Haber, Jemielniak, Kurasiński i Przegalińska, 2025).

Instytucjonalną ramę tego trendu wyznaczają dynamicznie powstające wytyczne, nadrzędne zasady jawności użycia AI, zachowania odpowiedzialności autorów oraz ciągłej aktualizacji reguł wraz z rozwojem technologii. Wytyczne tego typu formułowane przez zespoły naukowe i organizacje naukowe kładą nacisk na przejrzystość, nadzór człowieka i bezpieczeństwo procesu badawczego (Bockting, van Dis, van Rooij, Zuidema i Bollen, 2023).

Rekomendacje dla pisania z AI w czasopismach medycznych doprecyzowują dwie praktyki, które szybko stają się standardem, niewskazywanie narzędzi jako współautorów oraz precyzyjne ujawnianie sposobu użycia AI w metodach lub podziękowaniach, towarzyszy temu ostrzeżenie przed fabrykacją cytowań i błędami merytorycznymi (Cheng, Calhoun i Reedy, 2025).

W recenzowaniu i ewaluacji od przeglądu po oceny grantowe AI będzie nadal wspierać selekcję, streszczanie i sygnalizowanie ryzyk etycznych, ale bez zastępowania osądu eksperckiego. Zwraca się uwagę na możliwe uprzedzenia modeli, konieczność pełnej kontroli człowieka i jawne procedury korzystania z narzędzi w recenzji (Haber, Jemielniak, Kurasiński i Przegalińska, 2025).

W perspektywie krótkoterminowej należy oczekiwać dalszego rozwoju dynamicznych przeglądów i automatyzacji etapów technicznych z jednoczesnym wzrostem standardów raportowania od rejestru decyzji edycyjnych po treść promptów. Wkład koncepcyjny w takie praktyki, łączące automatyzację z rygiem przedstawiają między innymi prace o łączeniu mapowania nauki i standardów metodologicznych w zespołach interdyscyplinarnych (Tomczyk, Petrescu, Falco i Alola, 2024).

Kierunek średnioterminowy to rozproszone, współtworzone środowiska publikacyjne, w których autorzy korzystają z agentów wspierających pisanie, a instytucje wdrażają polityki „human-in-the-loop”. Zarówno literatura zarządcza, jak i popularnonaukowa opisuje te zmiany jako przejście od „efektywności” do „współkreacji” pod warunkiem zachowania odpowiedzialności, jawności i standardów oryginalności (Przegalińska i Jemielniak, 2023; Przegalińska i Triantoro, 2021).

Na horyzoncie widać też poszerzanie multimodalności (tekst-dane-obraz), integrację z otwartą nauką i wyższe wymagania dowodowe wobec modeli (wyjaśnialność, ścieżki dowodowe). Oczekiwana jest dalsza profesjonalizacja praktyk, obowiązkowe sprawdzanie oryginalności, krzyżowa weryfikacja cytowań w bazach naukowych, a także kształcenie w zakresie piśmienności AI po stronie autorów i recenzentów (Haber, Jemielniak, Kurasiński i Przegalińska, 2025).

Przyszłość pisania i recenzowania nie polega na automatyzacji osądu, lecz na lepszym wsparciu procesu naukowego. Wygrywać będą praktyki, które łączą automatyzację pracochłonnych etapów z rygorystyczną oceną ekspercką, przejrzystością i aktualizowanymi standardami. Dokładnie w tym kierunku zmierzają propozycje dynamicznych przeglądów i wytycznych oraz badania nad hybrydowymi metodykami syntezy (Tomczyk, Brüggemann i Vrontis, 2024; Beshr, Ateeq, Ateeq i Alaghbari, 2025).

Podsumowanie rozdziału 4

Omówione w rozdziale przykłady wskazują, że sztuczna inteligencja staje się nieodłącznym elementem nowoczesnej metodologii, wprowadzając dynamiczne przeglądy literatury, automatyzację wybranych etapów badań i nowe formy współpracy między autorami a recenzentami. Jej rola nie polega na zastępowaniu badacza, lecz na wzmacnianiu procesu twórczego i analitycznego, pod warunkiem zachowania odpowiedzialności i przejrzystości. Przyszłe kierunki rozwoju obejmują coraz szersze wykorzystanie multimodalnych modeli, integrację z otwartą nauką oraz rosnące wymagania dotyczące wyjaśnialności i transparentności działań AI.

6. Polityki wydawnictw naukowych i kwestie prawne

Upowszechnienie narzędzi generatywnej sztucznej inteligencji sprawiło, że wydawnictwa naukowe i instytucje akademickie zaczęły tworzyć szczegółowe regulacje dotyczące ich użycia. Polityki te odnoszą się do takich zagadnień jak autorstwo, transparentność, ochrona praw autorskich, etyka publikacyjna czy wykorzystanie AI w procesie recenzji. W niniejszym rozdziale przedstawiono stanowiska czołowych wydawnictw oraz kluczowe ramy prawne, które kształtują zasady korzystania z AI w nauce, wskazując jednocześnie na ich zróżnicowanie oraz wspólne elementy minimalne.

5.1. Stanowiska czołowych wydawnictw i czasopism naukowych

Polityki czołowych wydawnictw wobec generatywnej AI są zróżnicowane od podejścia zachowawczego, które ogranicza AI do korekty językowej, po bardziej otwarte, dopuszczające wsparcie w pisaniu czy analizie przy pełnej odpowiedzialności autora. Mimo różnic w detalach na przykład w progu ujawniania, wyjątkach dla badań nad samą AI, wyraźnie widać wspólny kanon: AI nie może być autorem, istotne użycie wymaga przejrzystego ujawnienia, a manipulowanie obrazami, danymi oraz powierzanie AI merytorycznej oceny manuskryptów jest zasadniczo zakazane. Polityki te dynamicznie się rozwijają, dlatego zawsze konieczna jest weryfikacja aktualnych wytycznych konkretnego wydawnictwa, czasopisma.

Elsevier dopuszcza użycie narzędzi generatywnej AI wyłącznie w celu poprawy czytelności i języka, pod warunkiem pełnej kontroli człowieka nad treścią i pełnej odpowiedzialności autorów za materiał, jak przedstawiono w tabeli 1. Wymagane jest jednoznaczne ujawnienie zastosowania AI w osobnej sekcji manuskryptu, jeśli narzędzie miało

udział w redakcji tekstu, przy czym AI nie może zostać wskazana jako autor lub współautor, ponieważ nie spełnia kryteriów odpowiedzialności za treść. W odniesieniu do materiału ilustracyjnego obowiązuje zakaz tworzenia lub modyfikowania obrazów naukowych z użyciem generatywnej AI; wyjątek dotyczy wyłącznie sytuacji, gdy jest to przedmiot badań i zostało przejrzystie opisane w metodach z możliwością żądania oryginalnych plików. W procesie recenzji zewnętrznej i redakcyjnej Elsevier zakazuje używania narzędzi typu ChatGPT do oceny prac, recenzja pozostaje w pełni „ludzka”, choć wydawca analizuje zgodne z polityką rozwiązania wspomagające.

Tabela 5.1. Zestawienie kluczowych zasad polityki Elsevier wobec generatywnej AI

Czy dozwolone użycie generatywnej AI w pisaniu tekstu?	Czy wymaga ujawnienia użycia AI?	Czy AI może być wymieniona jako autor/współautor?	Polityka dot. obrazów / figur generowanych przez AI	Polityka dot. recenzji (peer review) i/lub inne uwagi
Tak, ale tylko do poprawy czytelności i języka. Zastosowanie wyłącznie pod kontrolą człowieka (autor ponosi pełną odpowiedzialność).	Tak. Należy umieścić oświadczenie w osobnej sekcji, jeżeli użyto AI do pisania.	Nie. AI nie może być współautorem, ponieważ nie spełnia kryteriów odpowiedzialności za treść.	Zakaz użycia generatywnej AI do tworzenia lub modyfikowania obrazów naukowych. Jedynym wyjątkiem jest, gdy to część samego projektu badawczego – wtedy musi być opisane w metodach i można wymagać oryginalnych plików.	Recenzenci i redaktorzy: obowiązuje zakaz używania narzędzi typu ChatGPT do recenzowania; Elsevier pracuje nad oceną zgodnych z polityką rozwiązań AI, ale na razie recenzja ma być w pełni ludzka.

Źródło: opracowanie własne na podstawie zasad polityki Elsevier (data: 04.2025)

Springer Nature pozwala na wykorzystanie AI jako wsparcia językowego lub narzędzia analitycznego, przy zachowaniu pełnej odpowiedzialności autora za treść. Na podstawie zestawienia w tabeli 2 warto wskazać, że w przypadku bardziej zaawansowanych zastosowań np. generowania fragmentów tekstu konieczne jest ujawnienie w sekcji metodyka lub innej odpowiedniej części artykułu, zwykła korekta językowa nie zawsze wymaga deklaracji. AI nie może być przypisana jako autor. W zakresie ilustracji obowiązuje zakaz tworzenia generatywnych obrazów, chyba że wynika to bezpośrednio z projektu badawczego i przy zachowaniu weryfikowalności danych oraz poszanowania praw autorskich. W recenzji wydawca prosi, by recenzenci nie wgrywali manuskryptów do narzędzi AI z uwagi na poufność,

jeśli jakiekolwiek wsparcie AI miało miejsce, należy je ujawnić, aby zachować transparentność procesu.

Tabela 5.2. Zestawienie kluczowych zasad polityki Springer Nature wobec generatywnej AI

Czy dozwolone użycie generatywnej AI w pisaniu tekstu?	Czy wymagane ujawnienie użycia AI?	Czy AI może być wymienion a jako autor/wsp ółautor?	Polityka dot. obrazów / figur generowanych przez AI	Polityka dot. recenzji (peer review) i/lub inne uwagi
Tak, pod warunkiem, że AI narzędzia służą wsparciu językowemu (lub analizie danych), a tekst i odpowiedzialność wciąż należą do autora.	Tak, w przypadku bardziej zaawansowanego użycia LLM (np. generowania treści), należy to zgłosić w sekcji „Methods” lub innej odpowiedniej części. Zwykle narzędzia do korekty językowej nie wymagają ujawnienia.	Nie. Narzędzia AI nie spełniają kryteriów autorstwa; autor musi w pełni odpowiadać za treść.	Zabronione wykorzystanie generatywnych AI do tworzenia ilustracji. Jedynie w ściśle określonych wyjątkach (gdy to część badań AI, a dane są weryfikowalne i poszanowane prawa autorskie).	Recenzja: Springer Nature prosi recenzentów, by nie wgrywali manuskryptów do narzędzi AI (obawy o poufność i rzetelność). Jeśli recenzent w jakikolwiek sposób posiłkował się AI w ocenie badań, powinien to ujawnić.

Źródło: opracowanie własne na podstawie zasad polityki Springer Nature (data: 04.2025)

Na podstawie zestawienia w tabeli 3, można stwierdzić, że Taylor and Francis dopuszcza wykorzystanie AI jako wsparcia autora do generowania pomysłów, ułatwiania pisania i poprawy języka przy pełnej odpowiedzialności merytorycznej po stronie autorów. Wymagane jest jasne oświadczenie, jakie narzędzie zastosowano i w jakim celu w metodyce lub podziękowaniach AI nie może występować jako autor. Wydawca jednoznacznie zabrania tworzenia bądź manipulowania obrazami i danymi przeznaczonymi do publikacji (np. dodawanie lub usuwanie elementów), uznając to za naruszenie zasad. W recenzji redaktorzy i recenzenci nie mogą wgrywać materiałów do narzędzi AI (ryzyko naruszenia poufności i praw własności), dopuszczalne jest co najwyżej językowe „wygładzenie” recenzji z zachowaniem pełnej odpowiedzialności i poufności.

Tabela 5.3. Zestawienie kluczowych zasad polityki Taylor & Francis wobec generatywnej AI

Czy dozwolone użycie	Czy wymagane	Czy AI może być wymieniona	Polityka dot. obrazów / figur	Polityka dot. recenzji (peer review) i/lub inne uwagi
-------------------------	-----------------	-------------------------------	----------------------------------	--

generatywnej AI w pisaniu tekstu?	ujawnienie użycia AI?	jako autor/współautor?	generowanych przez AI	
Tak, ale ograniczone do wsparcia autora, np. w generowaniu pomysłów, ułatwianiu pisania, poprawie języka. Autor wciąż w pełni odpowiada za rzetelność i poprawność.	Tak. Konieczne jest jawne oświadczenie, które narzędzie i w jakim celu zostało wykorzystane (np. w sekcji Methods, Acknowledgments).	Nie. AI nie może być autorem. Autorzy muszą być ludźmi (odpowiedzialność prawna i merytoryczna).	Zabronione użycie AI do tworzenia i manipulowania obrazami, danymi do publikacji (np. usuwanie, dodawanie elementów). To narusza zasady T&F.	Recenzja: redaktorzy i recenzenci nie mogą wgrywać materiałów do narzędzi AI (ryzyko naruszenia własności intelektualnej i poufności). Recenzenci mogą ewentualnie korzystać z AI do poprawy języka recenzji, ale wciąż w pełni odpowiadają za treść i poufność.

Źródło: opracowanie własne na podstawie zasad polityki Taylor & Francis (data: 04.2025)

Wiley traktuje AI wyłącznie jako narzędzie wspomagające. Autorzy muszą zachować kontrolę nad całością tekstu i odpowiadają za poprawność merytoryczną oraz cytowania, co wskazano w tabeli 4. Ujawnienie użycia AI jest obowiązkowe w metodyce, podziękowaniach, gdy wpłynęło ono na treść poza prostą korektą pisowni i gramatyki. AI nie może być autorem, co Wiley dodatkowo osadza w wytycznych COPE (*Committee on Publication Ethics*, Komisja ds. Etyki Publikacji). W obrazach i danych zabronione jest generowanie lub manipulowanie wynikami, dopuszczalne są jedynie drobne korekty techniczne na przykład wyostrenie, poprawa barw, o ile nie zniekształcają danych. W recenzji możliwe jest użycie AI do usprawnienia języka samej recenzji, ale nie wolno wgrywać całych manuskryptów i należy ujawnić taki udział AI, aby chronić poufność.

Tabela 5.4. Zestawienie kluczowych zasad polityki Wiley wobec generatywnej AI

Czy dozwolone użycie generatywnej AI w pisaniu tekstu?	Czy wymagane ujawnienie użycia AI?	Czy AI może być wymienione jako autor /współautor?	Polityka dot. obrazów / figur generowanych przez AI	Polityka dot. recenzji (peer review) i/lub inne uwagi
Tak, ale wyłącznie jako narzędzie wspomagające. Autorzy muszą zachować	Tak. Koniecznie należy ujawnić w sekcji Methods/Acknowledgments	Nie. Wiley odwołuje się też do wytycznych COPE	Niedozwolone generowanie lub manipulowanie oryginalnymi danymi badawczymi czy wynikami przy użyciu AI.	Recenzja: Wiley sugeruje, że recenzenci/edytorzy mogą ewentualnie używać AI do usprawnienia języków

ć kontrolę i odpowiedzialność za całość tekstu, w tym za poprawność merytoryczną i cytowania.	ments, jeśli AI wpłynęła na treść (poza prostą korektą pisowni/gramatyki).	w tym zakresie. AI nie może posiadać praw autorskich ani wypełniać roli świadomego autora.	W odniesieniu do obrazów, Wiley nie pozwala na „kreowanie” wyników. Natomiast narzędzia do korekty (np. poprawa barw w ilustracjach) są dozwolone, o ile nie fałszują danych.	ego recenzji, ale nie wolno wgrywać całych manuskryptów do AI (ochrona poufności). Trzeba ujawnić, jeśli AI brała udział w tworzeniu recenzji.
---	--	--	---	--

Źródło: opracowanie własne na podstawie zasad polityki Wiley (data: 04.2025)

SAGE dopuszcza wykorzystywanie narzędzi do korekty stylu i języka co do zasady bez wymogu ujawnienia, jednak generatywne użycie AI tj. tworzenie faktycznych fragmentów publikacji musi zostać zadeklarowane wraz ze wskazaniem modelu i celu. AI nie może być przypisana jako autor, wydawca podkreśla obowiązek weryfikacji informacji, by unikać halucynacji i błędnych cytowań. W odniesieniu do ilustracji obowiązuje zakaz generowania obrazów przedstawiających dane, wyniki eksperymentów lub inne kluczowe elementy pracy; każdy użytek AI przy tworzeniu ilustracji należy ujawnić, przy czym redakcja może odrzucić materiał w razie wątpliwości co do rzetelności. W recenzji SAGE zabrania tworzenia recenzji z użyciem generatywnej AI z uwagi na poufność, dopuszczając jedynie narzędzia do edycji językowej.

Tabela 5.5. Zestawienie kluczowych zasad polityki SAGE wobec generatywnej AI

Czy dozwolone użycie generatywnej AI w pisaniu tekstu?	Czy wymagane ujawnienie użycia AI?	Czy AI może być wymieniona jako autor/współautor?	Polityka dot. obrazów / figur generowanych przez AI	Polityka dot. recenzji (peer review) i/lub inne uwagi
Tak, w zakresie wspomaganie pisania, np. narzędzia do korekty stylu i języka (bez konieczności ujawnienia). Natomiast korzystanie z generatywnych AI (np. ChatGPT) do tworzenia faktycznych fragmentów	Tak, w przypadku, gdy AI generowała fragmenty tekstu, dane lub inne treści. Trzeba wskazać, jaki model i w jakim celu został użyty.	Nie. AI nie może być przypisana jako autor. Należy zadbać o oryginalność treści i weryfikować informacje, które AI wygenerowała (np. unikając „halucynacji” i błędów w cytowaniach).	Zabronione generowanie obrazów, jeżeli służy to przedstawieniu danych, wyników eksperymentów czy innych kluczowych elementów publikacji. Należy ujawnić każde użycie narzędzi AI do tworzenia ilustracji – z zastrzeżeniem, że SA	Recenzja: SAGE zabrania wykorzystywać generatywne AI do tworzenia recenzji (naruszenie poufności). Dopuszcza użycie narzędzi do korekty językowej recenzji. Jeśli redaktor ma podejrzenia o niewłaściwe użycie A

publikacji musi być ujawnione.			GE może je odrzucić w razie wątpliwości co do rzetelności.	I (np. w tekście), może odrzucić zgłoszenie.
--------------------------------	--	--	--	--

Źródło: opracowanie własne na podstawie zasad polityki SAGE (data: 04.2025)

Jak wskazano w tabeli 6, IEEE (*Institute of Electrical and Electronics Engineers*) pozwala wykorzystywać AI w pisaniu pod warunkiem zachowania pełnej odpowiedzialności człowieka za treść i wymaga jawnego wskazania sekcji artykułu, w których użyto treści wygenerowanych przez AI, wzmianka powinna znaleźć się w podziękowaniach. AI nie może być autorem ani przejąć odpowiedzialności prawnej i etycznej. Analogicznej transparentności wymaga się w przypadku generowania obrazów lub kodów źródłowych, przy braku przyzwolenia na manipulowanie kluczowymi danymi. W recenzowaniu brak jest zgody na wgrywanie całych manuskryptów do narzędzi AI, jeśli AI posłużyła jedynie do poprawy języka recenzji, zaleca się to ujawnić z myślą o przejrzystości i ochronie poufności.

Tabela 5.6. Zestawienie kluczowych zasad polityki IEEE wobec generatywnej AI

Czy dozwolone użycie generatywnej AI w pisaniu tekstu?	Czy wymagane ujawnienie użycia AI?	Czy AI może być wymieniona jako autor/współautor?	Polityka dot. obrazów / figur generowanych przez AI	Polityka dot. recenzji (peer review) i/lub inne uwagi
Tak, można używać AI w pisaniu, ale konieczna jest pełna odpowiedzialność człowieka.	Tak. IEEE wymaga jawnego ujawnienia i wskazania sekcji artykułu, w których użyto treści generowanych przez AI. Wzmianka powinna znaleźć się w podziękowaniach (Acknowledgments).	Nie. AI nie może zostać uznana za autora, nie może też przejąć odpowiedzialności prawnej i etycznej.	Wymagane analogiczne ujawnienie, jeśli generowane były m.in. obrazy czy kody źródłowe. Nie ma otwartej zgody na manipulowanie kluczowymi danymi przy użyciu AI.	Recenzja: brak formalnego przyzwolenia na wgrywanie całych manuskryptów do narzędzi AI. Jeśli użyto AI do np. poprawy języka recenzji, zaleca się ujawnienie takiego faktu, aby zachować transparentność i poufność procesu.

Źródło: opracowanie własne na podstawie zasad polityki IEEE (data: 04.2025)

Emerald dopuszcza użycie AI do wsparcia językowego i poprawy tekstu, natomiast wszystkie nowe treści stworzone przez AI muszą być ujawnione, autorzy ponoszą pełną odpowiedzialność za rzetelność, co zestawiono w tabeli 7. AI nie może być autorem. W materiałach graficznych wydawca zabrania zgłaszania obrazów w całości wygenerowanych

przez AI, nie jest też wskazane generowanie obszernych fragmentów artykułu bez weryfikacji autorskiej. W recenzji redaktorzy i recenzenci nie mogą używać narzędzi generatywnych do przetwarzania artykułów i podejmowania decyzji wydawniczych, polityka Emerald wprost zakazuje tego rodzaju zastosowań w procesie peer review.

Tabela 5.7. Zestawienie kluczowych zasad polityki Emerald wobec generatywnej AI

Czy dozwolone użycie generatywnej AI w pisaniu tekstu?	Czy wymagane ujawnienie użycia AI?	Czy AI może być wymieniona jako autor/współautor?	Polityka dot. obrazów / figur generowanych przez AI	Polityka dot. recenzji (peer review) i/lub inne uwagi
Tak, ale ograniczone do wsparcia językowego i poprawy tekstu. Wszelkie nowe treści stworzone przez AI muszą zostać ujawnione. Autor w pełni odpowiada za rzetelność.	Tak, autor musi za deklarować użycie narzędzi generatywnych w artykule/rozdziale. Polityka obowiązuje od momentu ogłoszenia.	Nie. AI nie może być autorem (brak możliwości ponoszenia odpowiedzialności). Autorzy są zobowiązani do podawania źródeł i zachowania integralności tekstu.	Zabronione jest zgłaszanie do publikacji obrazów wygenerowanych w całości przez AI. Niewskazane jest też generowanie całych fragmentów artykułu, które nie są weryfikowane przez autora.	Recenzja: redaktorzy i recenzenci nie mogą używać narzędzi generatywnych do przetwarzania artykułów i decydowania o publikacji (poufność, ryzyko naruszenia praw). Emerald wprost zakazuje użycia AI w procesie recenzowania.

Źródło: opracowanie własne na podstawie zasad polityki Emerald (data: 04.2025)

W skali porównawczej Elsevier i Emerald należą do bardziej rygorystycznych, pozwalają na korzystanie z AI głównie w funkcjach korekty języka i stylistyki, podczas gdy Springer Nature, Taylor & Francis, Wiley, SAGE oraz IEEE dopuszczają nieco szersze wsparcie, na przykład wstępna analiza czy pomoc w redakcji, o ile autor bierze pełną odpowiedzialność i wprost ujawnia użycie AI w partiach twórczych. Wspólne minimum dla wszystkich to:

1. obowiązek transparentności: ujawnienie użycia AI zawsze wtedy, gdy narzędzie współtworzy treść lub ilustracje, prosta korekta językowa bywa z tego zwolniona na przykład SAGE, Springer;
2. jednoznaczna odmowa przyznania AI statusu autora - brak wyjątków;

3. ostrożność wobec obrazów i danych: powszechny zakaz generowania lub modyfikowania materiału, który mógłby zafałszować wyniki, Elsevier i Springer wyjątkowo stanowczo zabraniają generowania obrazów naukowych poza jasno opisanym kontekstem badań nad samą AI, Wiley dopuszcza jedynie drobne korekty techniczne;
4. rygory w recenzji: brak zgody na wykorzystanie AI do merytorycznej oceny prac i wgrywanie manuskryptów do narzędzi AI, przy ewentualnej, ograniczonej akceptacji prostych poprawek językowych recenzji z obowiązkiem ujawnienia.

W odniesieniu do polskich czasopism i wydawnictw wciąż bywa, że brak im formalnych, szczegółowych wytycznych dotyczących AI, praktycznym minimum pozostaje merytoryczne uzasadnienie zastosowania oraz pełna transparentność wobec redakcji i czytelników.

Coraz częściej pojawiają się jednak wyjątki wyznaczające klarowny standard, na przykład periodyk *Journal of Modern Science* przyjął kompleksową politykę AI, wymaga precyzyjnego ujawnienia (nazwa, model narzędzia, wersja i data użycia, treść promptów), dopuszcza AI wyłącznie do korekty językowej, stylistyki i tłumaczeń, a zakazuje generowania treści merytorycznej i danych, AI nie może być autorem.

W zakresie ilustracji zabronione jest tworzenie lub modyfikowanie obrazów przy użyciu AI z wyjątkiem sytuacji, gdy jest to element metod badawczych, szczegółowo opisany, dopuszcza się jedynie generowanie wykresów na podstawie danych przygotowanych przez autora, traktując AI jak narzędzie graficzne, z obowiązkiem odnotowania tego w artykule.

Polityka obejmuje też recenzentów i redaktorów: bezwzględna poufność materiałów i korespondencji, zakaz wprowadzania manuskryptów i raportów do narzędzi AI oraz wyłączna odpowiedzialność człowieka za oceny i decyzje; nieujawnione użycie AI grozi odrzuceniem pracy, a redakcja zastrzega możliwość detekcji AI w Plagiat.pl (*Journal of Modern Science*).

Takie rozwiązania zbliżają praktykę krajową do globalnego kanonu przejrzystości i odpowiedzialności oraz mogą służyć jako wzorzec przy kształtowaniu lokalnych polityk.

Łącznie polityki budują spójny standard - AI może wspierać autorów jako edytor języka i narzędzie pomocnicze, ale nie może zastępować autora, nie może kreować lub poprawiać wyników, a każdy ślad jej udziału w treściach wykraczających poza korektę powinien być jawny i kontrolowalny.

5.2. Kryteria autorstwa a użycie AI; deklaracje transparentności

Kwestia autorstwa w nauce nie sprowadza się wyłącznie do formalnego przypisania nazwisk twórcy danej publikacji naukowej. Od zawsze związana jest z odpowiedzialnością za

treść pracy, integralność oraz otwartością do podjęcia krytycznej dyskusji nad przedstawionymi wynikami. W kontekście pojawienia się narzędzi generatywnej sztucznej inteligencji otworzyły się nowe pola problemów dotyczącego tego obszaru. Zasadnicze pytanie dotyczy bowiem tego, jak rozumieć wkład technologii w proces badawczy i jak unikać przypisywania jej cech, które tradycyjnie były zarezerwowane wyłącznie dla autorów.

Zgodnie ze stanowiskiem Committee on Publication Ethics (COPE) sztuczna inteligencja nie może być uznana za autora, gdyż nie spełnia podstawowych kryteriów, takich jak zdolność do podejmowania świadomych decyzji, ponoszenia odpowiedzialności czy reagowania na proces recenzyjny (COPE, 2023). Podobnie w licencjach największych dostawców narzędzi genAI znajdziemy zapisy podkreślające brak możliwości przypisania autorstwa na poczet technologii.

Tabela 5.8. Zestawienie do praw autorskich treści wygenerowanych z wykorzystaniem genAI

Firma / Narzędzie	Status autorstwa treści wygenerowanych przez AI	Kluczowe zapisy / praktyki
OpenAI – ChatGPT	Autorstwo treści przypisane użytkownikowi. OpenAI podkreśla, że nie rości sobie praw do Outputu.	Użytkownik zachowuje prawa do Inputu i Outputu; OpenAI może wykorzystywać dane do szkolenia, o ile użytkownik nie wyłączy tego w ustawieniach. (OpenAI, 2025)
Microsoft – Copilot	Autorstwo treści przypisane klientowi/użytkownikowi. Microsoft deklaruje, że nie przejmuje praw autorskich.	Microsoft zobowiązuje się bronić użytkowników przed roszczeniami o naruszenie praw autorskich, jeśli korzystają zgodnie z zaleceniami i filtrami. (Microsoft, 2023)
Google – Gemini	Autorstwo traktowane jako należące do użytkownika, przy jednoczesnym silnym nacisku na transparentność (wskazanie, że treść pochodzi z AI).	Użytkownik zachowuje prawa do treści, ale Google zastrzega możliwość wykorzystania ich do udoskonalania usług; stosuje narzędzia znakowania. (Google, 2024)
Anthropic – Claude	Autorstwo treści przypisane użytkownikowi. Claude pełni rolę narzędzia wspierającego, nie podmiotu praw autorskich.	Użytkownik posiada prawa do danych wyjściowych, ale Anthropic zastrzega, że nie ponosi odpowiedzialności za ewentualne roszczenia osób trzecich wynikające z naruszenia praw. (Anthropic, 2023)

Źródło: opracowanie własne na podstawie zasad polityk twórców modeli OpenAI, Google, Microsoft, Anthropic

Kiedy przyjrzymy się dla przykładu jednemu z głównych graczy na rynku genAI, jakim jest OpenAI, zauważymy, że kwestia autorstwa jest uregulowana w sposób stosunkowo przejrzysty. Zgodnie z obowiązującymi zasadami, użytkownik zachowuje pełne prawa do treści,

które wprowadza do systemu (*input*) oraz do treści wygenerowanych przez model (*output*). OpenAI wprost deklaruje, że nie rości sobie praw autorskich do efektów pracy użytkowników, zastrzegając jedynie konieczność respektowania ograniczeń wynikających z obowiązującego prawa (np. ochrony danych osobowych czy dóbr osobistych) oraz własnych regulacji wewnętrznych. Z tego względu, niezależnie od wybranego narzędzia, dla młodego naukowca ważna jest interpretacji z perspektywy prawa autorskiego, wytycznych macierzystych jednostek naukowych, wydawnictw oraz etyki naukowej.

Ważnym aspektem w zakresie praw autorskich jest też kwestia udostępniania własnych zasobów lub utworów (w rozumieniu prawa autorskiego) np. w formie spersonalizowanych GPT-ów. Jeżeli zostają one udostępnione publicznie w GPT Store autor treści udziela OpenAI niewyłącznej, ogólnoswiatowej, nieodwołalnej i wolnej od opłat licencji na korzystanie, przechowywanie, modyfikowanie, dystrybucję i promowanie stworzonych materiałów. Oznacza to, że akt kreacji w środowisku OpenAI możemy podzielić na prywatny i publiczny. W sferze prywatne regulacje zapewniają o prywatności prowadzonych konwersacji i projektów.

Z perspektywy kryterium autorstwa kluczowe jest rozróżnienie pomiędzy aktem twórczym człowieka a działaniem modelu. OpenAI podkreśla, że choć użytkownik nabywa prawa autorskie do wygenerowanego materiału, to nie może być on uznany za jedynego autora w sensie tradycyjnym. Algorytm pełni bowiem rolę narzędzia, które współuczestniczy w procesie twórczym, który pomimo tego faktu nie może być podmiotem prawa autorskiego. W praktyce oznacza to, że kryterium autorstwa sprowadza się do stopnia oryginalnego wkładu człowieka w przygotowanie promptu, redakcję i interpretację uzyskanych treści. Autorstwo w środowisku AI ma więc charakter procesualny i relacyjny, a nie indywidualistyczny i zamknięty.

W związku z tym narzędzie genAI stanowi funkcję wspierającą, a jego wykorzystanie powinno zostać jasno zadeklarowane. Tak rozumiana transparentność obejmuje wskazanie zakresu, celu i charakteru użycia. Należy wskazać, czy dotyczyło to edycji językowej, wsparcia analitycznego czy też generowania wizualizacji. Ten obowiązek transparentności znajduje swoje umocowanie również w dokumentach dotyczących kwestii badań naukowych w Polsce. Kodeks etyki pracownika naukowego PAN (2024) w paragrafie 16 jednoznacznie wskazuje, że każdy wkład w pracę naukową powinien być rzetelnie i transparentnie opisany, a odpowiedzialność za całość publikacji ponoszą wyłącznie jej autorzy. W świetle tych zapisów, zgodnie z wcześniej podanymi zapisami i interpretacjami, nie ma miejsca na przypisywanie

autorstwa technologii, która nie posiadają sprawczości ani zdolności moralnej odpowiedzialności (PAN, 2024).

Podobne stanowisko zajmują krajowe inicjatywy związane z edukacją i rozwojem kompetencji oraz korzystaniem z cyfrowych technologii. Wytyczne przygotowane przez zespół Centrum Nowoczesnej Edukacji Politechniki Gdańskiej kierowany przez prof. Joannę Mytnik, podkreślają, że narzędzia generatywne mogą wspierać proces naukowy i dydaktyczny, jednak ich użycie powinno być ujawniane w taki sposób, aby umożliwiała ocenę faktycznego wkładu autora. Transparentność staje się gwarancją nie tylko rzetelności badań, ale również przejrzystości w kształceniu młodych naukowców (CNE, 2023). Wobec

zachodzących zmian warto podkreślić, że przyszłość pracy naukowej wymaga równowagi między korzystaniem z potencjału AI a zachowaniem centralnej roli badacza. Autorstwo jest nierozdzielnie związane z odpowiedzialnością, a deklaracje transparentności pełnią funkcję ochronną, wzmacniającą zaufanie i integralność środowiska akademickiego. Dla młodych naukowców oznacza to konieczność budowania warsztatu badawczego opartego na świadomości etycznej i refleksyjnego budowania synergii człowieka i technologii.

5.3. Prawa autorskie do tekstu oraz modele licencjonowania

Prawo autorskie, nie tylko w obszarze pracy naukowej jest złożonym i wielopoziomowym zagadnieniem, którego wyczerpanie nawet w obszernej publikacji jest trudne do uzyskania. Dlatego w niniejszym podrozdziale znajduje się temat otwartych licencji, który powstał przed upowszechnieniem modeli genAI, a jest uniwersalny i warto dalszego rozwoju i wspierania zarówno ze względu na czytelność, powszechność, jak i kwestie etyczne.

Warto zauważyć, że problematyka autorstwa w świecie cyfrowym nie ogranicza się wyłącznie do regulacji wewnętrznych dostawców usług genAI. Od początku XXI wieku istotną rolę w porządkowaniu i ujednolicaniu obszaru otwartych zasobów odgrywały licencje Creative Commons (CC), które stanowiły praktyczną odpowiedź na rosnącą potrzebę klarownego i międzynarodowo zrozumiałego określania warunków korzystania z twórczości¹. Dzięki nim możliwe stało się stworzenie globalnej infrastruktury zaufania – systemu, w którym twórcy mogli świadomie dzielić się swoją pracą, a odbiorcy mieli pewność co do zakresu dopuszczalnego wykorzystania.

Współczesne praktyki publikacyjne coraz częściej łączą rygory tradycyjnych modeli ochrony prawnej z ideą otwartego dostępu, której kluczowym narzędziem stały się licencje Creative Commons (CC). Ich znaczenie w nauce polega na tym, że pozwalają autorom

zachować prawa osobiste, jednocześnie umożliwiając różny stopień swobodnego wykorzystania utworu przez innych badaczy i instytucje.

System licencji CC obejmuje kilka wariantów, które różnią się zakresem przyznanych odbiorcom uprawnień. Licencja CC BY dopuszcza dowolne wykorzystanie utworu pod warunkiem podania autorstwa i należy do najczęściej stosowanych w publikacjach naukowych, ponieważ maksymalizuje potencjał rozpowszechniania i cytowania. Wersja CC BY-SA nakłada dodatkowy obowiązek dzielenia się dziełami pochodnymi na tych samych zasadach, co wzmacnia kulturę współpracy i dzielenia się wiedzą. Ograniczenia pojawiają się w wariantach takich jak CC BY-NC, które zakazują użycia komercyjnego, czy CC BY-ND, które uniemożliwiają tworzenie opracowań, a więc zmniejszają elastyczność wykorzystania w kontekście badań i dydaktyki.

5.4. Dostosowanie narzędzi licencyjnych

Pojawienie się generatywnej sztucznej inteligencji w znaczący sposób skomplikowało ten porządek. Modele językowe były trenowane na ogromnych zbiorach danych, w których kwestie autorstwa i warunki licencyjne pozostawały w dużej mierze pominięte, szczególnie na wczesnym etapie rozwoju technologii. Oznacza to, że fundament otwartości, wzajemność, poszanowanie praw i przejrzystość, został w praktyce zawieszony. W tym kontekście Creative Commons, jako ruch budujący etyczne ramy dla otwartej kultury, może stanowić punkt odniesienia i źródło rozwiązań dla nowych wyzwań.

Jedną z odpowiedzi na to wyzwanie jest inicjatywa CC Signals, zaprezentowana w 2025 roku jako próba stworzenia „nowego kontraktu społecznego na erę AI” (Creative Commons, 2025). Signals mają formę społeczno-technicznego standardu, który pozwala twórcom i instytucjom wyrażać preferencje dotyczące wykorzystania ich treści przez systemy sztucznej inteligencji, zarówno w procesie treningu, jak i w dalszej dystrybucji. Zasadniczym celem Signals jest zachowanie ducha Creative Commons: prostoty, dostępności i globalnej użyteczności, przy jednoczesnym dostosowaniu do realiów uczenia maszynowego. Co istotne, rozwiązanie to nie odrzuca dorobku wcześniejszych licencji, lecz go rozwija, budując most pomiędzy otwartością a odpowiedzialnym wykorzystaniem danych w epoce GenAI.

W ten sposób Creative Commons może nadal pełnić rolę etycznego regulatora przestrzeni otwartych zasobów, tym razem obejmując także sztuczną inteligencję. CC Signals są propozycją, która zachowuje kluczowe wartości: transparentność, wzajemność i poszanowanie twórczości. Jednocześnie rozszerza je na dynamicznie zmieniające się środowisko technologiczne. Właśnie dlatego stanowią one przykład praktyki, która, opierając

się na wypracowanych i sprawdzonych rozwiązaniach, odpowiada na globalne potrzeby i pozostaje dostępna dla całej społeczności międzynarodowej.

Obok wolnych licencji funkcjonują także inne modele licencjonowania stosowane w publikacjach naukowych. Wydawnictwa komercyjne często wymagają przekazania pełnych praw majątkowych na ich rzecz, co prowadzi do sytuacji, w której autorzy tracą kontrolę nad sposobem udostępniania własnych tekstów. Coraz częściej alternatywą stają się modele hybrydowe, w ramach których autorzy mogą wybrać między klasycznym zamkniętym dostępem, a publikacją w trybie open access, zazwyczaj jednak wymagającym opłaty. Istotne są także licencje wydawnicze oparte na modelu Gold Open Access, które zapewniają pełne otwarcie zasobów na zasadzie CC BY, oraz modelu Green Open Access, umożliwiającego autoarchiwizację wersji manuskryptu w repozytoriach instytucjonalnych.

Z punktu widzenia młodych naukowców kluczowe znaczenie ma zrozumienie, że wybór licencji nie jest wyłącznie kwestią techniczną, lecz decyduje o tym, jak dalece ich prace mogą uczestniczyć w obiegu wiedzy i budować zbiór zasobów naukowych. Licencje otwarte zwiększają widoczność publikacji, sprzyjają współpracy międzynarodowej i ułatwiają budowanie synergii między badaczami. Jednocześnie wymagają świadomości prawnych ograniczeń i odpowiedzialności za to, aby udostępniane treści nie naruszały cudzych praw.

Podsumowanie rozdziału 5

Analiza polityk wydawniczych i regulacji prawnych pokazuje, że choć poszczególne instytucje różnią się w podejściu do wykorzystania AI, wszystkie podkreślają konieczność przejrzystości, zakaz uznawania modeli za autorów oraz obowiązek odpowiedzialności po stronie badacza. Coraz większe znaczenie mają także kwestie praw autorskich, ochrony danych i zgodności z regulacjami takimi jak RODO czy AI Act. Przyszłość praktyki publikacyjnej będzie opierać się na łączeniu otwartości na innowacje z rygorystycznymi zasadami etycznymi i prawnymi, aby zapewnić integralność procesu naukowego.

Część II przewodnik praktyczny

Druga część monografii koncentruje się na praktycznym wymiarze wykorzystania sztucznej inteligencji w pracy naukowej. W części teoretycznej przedstawiono podstawowe pojęcia, definicje i konteksty etyczno-prawne, natomiast tutaj uwaga skupia się na konkretnych narzędziach, technikach i scenariuszach ich zastosowania. Zaprezentowane rozwiązania obejmują zarówno wsparcie w pisaniu i redakcji tekstów, jak i automatyzację zadań pomocniczych, analizę literatury czy weryfikację treści generowanych przez modele językowe. Celem tej części jest dostarczenie doktorantom i młodym badaczom zestawu praktycznych wskazówek, które pozwolą świadomie i odpowiedzialnie włączać narzędzia AI do własnego warsztatu badawczego, z zachowaniem krytycznej refleksji i kontroli jakości rezultatów.

7. Inżynieria promptów (prompt engineering)

Niniejszy rozdział skupi się na praktycznych aspektach korzystania z generatywnej sztucznej inteligencji w pracy młodego naukowca. Mając w pamięci definicje, informacje i punkt widzenia zawarty w pierwszej części zachęcamy do praktycznego eksplorowania narzędzi AI oraz poszukiwania własnej drogi i rozwijania kompetencji.

Wśród wielu narzędzi z jakich możemy skorzystać znajdziemy na pewno chatGPT, Gemini, Copilota, Claude, Notebooka LM, czy polskie modele w postaci Bielika, czy PLLuM. Pomimo tego, że lista narzędzi jest o wiele dłuższa, i z każdym dniem pojawiają się nowe rozwiązania, czy też kolejne wersje dostępnych narzędzi, specyfika i techniki pracy z nimi jest uniwersalna. Oznacza to, że możemy rozwijać swoje doświadczenia pracując z jednym narzędziem, a następnie wykorzystywać zdobyte kompetencje w pracy w innych środowiskach.

Sposób pracy z genAI pozwalający na generowanie wartościowych treści poprzez pisanie odpowiednio skonstruowanych poleceń w języku naturalnym to inżynieria promptów. „Inżynieria promptów (Prompt engineering) to niezbędna technika, która rozszerza możliwości dużych modeli językowych (LLM) oraz modeli wizualno-językowych (VLM). Polega ona na strategicznym projektowaniu instrukcji specyficznych dla danego zadania, nazywanych promptami, w celu kierowania wynikiem modelu bez modyfikowania jego podstawowych parametrów.” (Sahoo i in., 2024)

Główne aspekty inżynierii promptów obejmują (Sahoo i in., 2024):

- zwiększanie skuteczności modelu – prompty pozwalają na bezproblemową integrację wstępnie wytrenowanych modeli z nowymi zadaniami. Wywołując pożądane zachowania modelu wyłącznie na podstawie podanej instrukcji.
- rodzaje promptów – mogą to być instrukcje w języku naturalnym, które dostarczają kontekst do ukierunkowania modelu, lub wyuczone reprezentacje wektorowe, które aktywują odpowiednią wiedzę.
- adaptowalność – oferuje mechanizm do precyzyjnego dostrajania wyników modelu za pomocą starannie opracowanych instrukcji, co pozwala modelom wykazywać się w różnorodnych zadaniach i dziedzinach. Ta elastyczność odróżnia ją od tradycyjnych metod. Te często wymagają ponownego szkolenia lub obszernego dostrajania parametrów modelu.
- transformacyjny wpływ – inżynieria promptów stała się transformacyjną siłą w dziedzinie sztucznej inteligencji, odblokowując ogromny potencjał LLM i umożliwiając im wykonywanie zadań, które wcześniej wydawały się niemożliwe, od generowania języka i odpowiadania na pytania po tworzenie kodu i zadania rozumowania.

Istotnym aspektem, który wskazuje jak bardzo zróżnicowana i spersonalizowana może być praca z wykorzystaniem promptowania, jest fakt, że nie ma jednego uniwersalnego modelu prompta (instrukcji), którą możemy stosować, aby uzyskać wartościowy i adekwatny do naszych potrzeb i sytuacji efekt. „Krajobraz współczesnej inżynierii promptów obejmuje spektrum technik, od metod fundamentalnych, takich jak zero-shot i few-shot prompting, po bardziej skomplikowane podejścia, takie jak 'chain of code' prompting.” (Fagbohun i in., 2024).

Szeroki zakres zagadnienia wymaga od młodych badaczy ustawicznego uczenia się, eksperymentowania, dociekliwości i korzystania z zaplecza własnych kompetencji, co zostało wspomniane w rozdziale pierwszym. Wtedy techniki inżynieringu promptowania mogą się integrować z całym warsztatem pracy dydaktycznej i naukowej.

6.1. Prompt engineering, a prompt learning

W literaturze przedmiotu coraz częściej dokonuje się rozróżnienia pomiędzy prompt engineering a prompt learning. Pierwsze z tych pojęć odnosi się do praktyki ręcznego formułowania poleceń w języku naturalnym, których celem jest uzyskanie od modelu odpowiedzi możliwie precyzyjnej i zgodnej z intencją użytkownika. To właśnie ta forma promptowania stanowi podstawowe narzędzie w pracy młodego badacza, pozwalając w sposób iteracyjny udoskonalać instrukcje i kontrolować jakość generowanych rezultatów. Jednak rozwój badań nad dużymi modelami językowymi pokazał, że nie wszystkie aspekty interakcji z nimi muszą odbywać się wprost poprzez język naturalny. Pojawił się nurt określany mianem

prompt learning, który obejmuje różne metody automatycznego uczenia samych promptów. Kluczowe znaczenie mają tu takie rozwiązania jak soft prompts czy prefix-tuning. W przypadku soft prompts zamiast tradycyjnych słów stosuje się ciągi wektorów w przestrzeni ukrytej modelu, które nie odpowiadają żadnym rzeczywistym jednostkom językowym, lecz pełnią rolę subtelnych sygnałów kierujących generowaniem odpowiedzi (Liu i in., 2021). Prefix-tuning natomiast polega na doklejaniu do wejścia modelu wyuczonych prefiksów wektorowych, które pozwalają wprowadzać kontekst zadania i sterować zachowaniem modelu bez konieczności modyfikacji wszystkich jego parametrów (Li i Liang, 2021).

Z punktu widzenia praktyki naukowej różnica pomiędzy tymi dwoma podejściami jest istotna. Prompt engineering jest narzędziem dostępnym dla każdego użytkownika i pozwala rozwijać umiejętność formułowania precyzyjnych poleceń. Prompt learning to z kolei obszar badań nad tym, jak modele mogą być dostosowywane do specyficznych zadań badawczych poprzez uczenie niewielkich fragmentów parametrów zamiast pełnego trenowania modelu. Zrozumienie obu perspektyw umożliwi młodemu badaczowi nie tylko świadome korzystanie z narzędzi dostępnych w pracy codziennej, ale także orientację w kierunkach badań nad nowymi sposobami interakcji z generatywną sztuczną inteligencją.

6.2. Struktura skutecznego procesu promptowania

Pomimo tego, że prompty mogą przyjmować różnorodne formy, istnieje schemat, którym może pomóc w poruszaniu się w świecie genAI. W ramach kursu „Google Prompting Essentials”, który dostępny jest m.in. na platformie Coursera, inżynierowie Googla dzielą się następującą strukturą procesu promptowania, którą przedstawiono w tabeli 6.1.

Tabela 6.1. Składowe procesu promptowania

Zadanie	Kontekst	Odniesienie	Ewaluacja	Powtórzenie
Szczegółowy opis celu, formy zadania oraz roli sztucznej inteligencji	Istotne informacje prezentujące zakres, analizę, specyfikę danego zagadnienia i potrzeb.	Przykłady, raporty, dane związane z zadaniem. Tło kulturowe, społeczne, paradygmat.	Ocena wygenerowanej treści	Kontynuacja procesu poprzez uzupełnienie luków, doprecyzowanie, uzupełnienie danych.

Źródło: opracowanie własne na podstawie Google Prompting Essentials

Podobne składowe promptowania znajdziemy w literaturze zajmującej się tym tematem. W artykule White’a znajdziemy następujące składowe (White, i in., 2023), które przedstawiono w tabeli 13.

Tabela 6.2. Składowe promptowania wg White i inni

Nazwa i klasyfikacja	Zamiar i kontekst	Motywacja	Struktura i kluczowe idee	Przykładowa implementacja	Konsekwencje
Identyfikuje i wskazuje problem do rozwiązania	Opisuje problem i cele, które mają być osiągnięte	Uzasadnienie problemu/zadania	Fundamentalne informacje kontekstowe	Sprawdzenie zastosowania promptu w praktyce	Refleksja na temat wad i zalet prompta

Źródło: opracowanie własne na podstawie White i inni (White, i in., 2023)

W dostępnych tekstach podkreśla się także znaczenie iteracji. Jest ona związana z ciągłym doskonaleniem promptów, co wskazuje na procesowy charakter komunikacji z narzędziami genAI (White, i in., 2023; Liu, i in., 2023, Fagbohun i in., 2024). Ponieważ model przygotowany przez inżynierów Google jest przejrzysty i zgodny z opracowaniami naukowymi prześledźmy poszczególne etapy procesu promptowania omawiając jego poszczególne punkty.

Zadanie

Przystępując do procesu wykorzystania dowolnego narzędzia LLM należy rozpocząć podobnie jak w pracy projektowej, od określenia celu, któremu ma służyć nasze działanie. Cel powinien być specyficzny i określać jasno do czego chcemy dążyć. Zupełnie inaczej będzie wyglądał proces przygotowywania instrukcji i przebiegu konwersacji w przypadku potrzeby uzyskania informacji zwrotnej, korekty językowej, przeglądu materiału, wsparcia w analizie danych czy tworzenia symulacji lub projektowania eksperymentu.

Na tym etapie możemy też określić jaką rolę chatbot powinien przyjąć. Może przecież sprawdzić czytelność naszego tekstu z punktu widzenia początkującego w danej materii studenta, ale może też udzielić na wskazówek jako konkretny ekspert w danej dziedzinie. Określenie jego charakterystyki wpłynie na to, co zostanie wygenerowane w czasie rozmowy. Istnieje także możliwość przypisania formy w zakresie długości tekstu, ograniczenia jego zakresu co do liczby słów, zdań, formy z punktowaniem lub ciągłej. Istnieje możliwość dookreślenia stylu pod względem formalnym, nieformalnym, popularnonaukowym czy akademickim.

Kontekst

Narzędzia generatywnej sztucznej inteligencji są maszynami cyfrowymi, które bazują na danych, do których mają dostęp. Nie charakteryzują się domyślanymi naszymi intencjami, potrzebami, paradygmatami. Kiedy otwieramy okno dialogowe i zadajemy pytanie, chatbot

odpowiada stosując zaawansowaną statystykę, korzystając z danych, na podstawie których został wyszkolony. Krótkie polecenie powoduje, że treść, którą nam dostarcza opiera na danych dostarczonych w korpusach treningowych (w tym pozyskanych z sieci), a jego odpowiedź jest wypadkową skonstruowaną dla potrzeb przeciętnego użytkownika.

Z całą pewnością takiej sytuacji chcemy uniknąć. Dlatego rozszerzając instrukcję o konkretne okoliczności i tło w ramach, którego odpowiedzi mają być generowane, powodujemy, że manipulujemy zachowaniem modelu (Liu, 2023). Dookreślenie zapytania poprzez udostępnienie naszych danych, które wskazują pożądaną przez nas zakres odpowiedzi, wskazania naszej roli oraz wymaganej perspektywy doprecyzowuje styl i zakres odpowiedzi generowanej przez genAI.

Poprzez wskazówki i ukierunkowanie na konkretne obszary dydaktyczne i badawcze możemy w taki sposób dostroić odpowiedzi modelu, aby były one bardziej satysfakcjonujące i dostosowane do naszych bieżących potrzeb.

Odniesienia

Wielu z nas posiada opracowania, wyniki badań, cenne dla niego teksty, zarówno własne, jak i te które są dla nas ważnym odniesieniem w rozwoju i pracy dydaktycznej i naukowej. To zasoby, które są nie do przecenienia. Ich wykorzystania wspiera i pogłębia wartość dostarczonego kontekstu. Załączenie ich w oknie dialogowym pozwala zawęzić obszar odpowiedzi do danego punktu widzenia, obszaru wiedzy, stylu analizy, czy syntezy.

Dodatkowo przygotowując instrukcję możemy także dostarczyć wzorzec odpowiedzi, w taki sposób, aby w konwersacji chatbot zapoznał się z jego strukturą i językiem, a następnie kierował się nim w swoich odpowiedziach. To, pozwala dostosować formę i styl generowanych treści do potrzeb użytkownika.

Nie bez znaczenia są tutaj także konkretne bazy danych, na których chcemy oprzeć własną pracę. Dostarczanie i wskazywanie, a także sprawdzanie czy model korzysta z tych baz w poprawny sposób, może podnieść jakość wyników konwersacji.

Ocena i powtórzenie

Wiele razy na kartach tej publikacji wspominaliśmy, że ludzkie kompetencje są niezbędne w procesie współpracy z cyfrowymi maszynami. Ta składowa procesu promptowania jest istotnym wkładem w jakość uzyskiwaną w ramach tego procesu. Z wygenerowanymi treściami należy się zapoznać, poddać analizie i krytycznemu osądowi. Stanowi to podstawę zarówno dla dalszego wykorzystania tego materiału w dalszych etapach pracy, ale także jest

ważnym elementem pozwalającym na to, aby w wartościowy sposób przejść do kolejnego cyklu promptowania, decydując się na dalszą współpracę w ramach danego zadania.

Powtarzanie i pogłębianie procesu promptowania jest niezbędnym elementem profesjonalnego korzystania z modeli. Nawet w najlepiej zaprojektowanym pierwszym etapie promptowania, pojawiają się błędy, nieporozumienia i luki kontekstowe, które wpływają na uzyskany efekt. Dlatego ich dostrzeżenie w procesie ewaluacji oraz ponowne dookreślanie potrzeb, zakresów, stylów i oczekiwań jest niezbędne.

Powyższe składowe procesu promptowania wskazują, że wartościowa praca z tymi narzędziami również wymaga pracy, zaangażowania i korzystania z własnych zasobów poznawczych.

6.3. Techniki prompt engineeringu

Praktyka pracy z dużymi modelami językowymi ujawnia, że sposób formułowania poleceń w znaczący sposób wpływa na jakość generowanych treści. W literaturze wyróżnia się kilka podstawowych technik, które pozwalają lepiej zrozumieć mechanizmy działania modeli i uzyskać rezultaty bardziej adekwatne do potrzeb badacza.

Jedną z najprostszych metod jest tzw. *zero-shot prompting*, polegające na zadawaniu pytania bez wcześniejszego podawania przykładu. Badacz korzysta tu z założenia, że model potrafi rozpoznać intencję wyłącznie na podstawie treści zapytania. W praktyce jednak, aby zwiększyć precyzję odpowiedzi, stosuje się *few-shot prompting*. Polega ono na poprzedzeniu właściwego pytania kilkoma przykładami oczekiwanych odpowiedzi, co pozwala modelowi uchwycić wzorzec i dostosować dalsze generowanie. Jeszcze innym podejściem jest *chain-of-thought prompting*, które zachęca model do prezentowania przebiegu rozumowania krok po kroku. Ta technika okazuje się szczególnie użyteczna w kontekstach akademickich, gdzie liczy się nie tylko wynik końcowy, ale także przejrzystość i logika wyводу (Wei i in., 2022).

Na uwagę zasługuje także tzw. *role prompting*, w którym modelowi przypisuje się określoną rolę, np. recenzenta czasopisma naukowego, promotora czy wykładowcy. Odpowiedzi wygenerowane w takim trybie odzwierciedlają nie tylko treść pytania, lecz także kontekst społeczny lub instytucjonalny, w jakim zostały umieszczone. Z kolei podejście *instruction-based prompting* opiera się na precyzyjnym sformułowaniu polecenia, często z wyraźnym wskazaniem oczekiwanej formy odpowiedzi, np. w postaci tabeli, listy czy streszczenia o określonej długości.

Każda z wymienionych technik pokazuje, że prompt engineering nie sprowadza się do intuicyjnego zadawania pytań, lecz wymaga świadomego wyboru strategii odpowiadającej celowi badawczemu. Umiejętność ich stosowania staje się elementem warsztatu naukowego, podobnym do znajomości metod analizy danych czy kryteriów poprawnego wnioskowania.

6.4. Sposoby budowania promptów

Jak wskazaliśmy we wcześniejszej części dotyczące procesu promptowania, struktura prompta może być zróżnicowana. Poniżej przedstawione zostały dwa podejścia do budowania promptów. Struktura prompta zaproponowana przez Macieja Chrzanowskiego, może wyglądać następująco (Wojewodzic, 2024, s.22):

Tablica 6.3. Przykładowa struktura prompta

Rola	Zadanie	Struktura zadania	Kontekst	Cel	Format odpowiedzi
Zachowuj się jak [...]	Opisz [...] w stylu [...]	Najpierw [...], a następnie [...] /Maksymalną a liczba znaków [...]	Opis ma być elementem [...]	Twoim celem będzie [...] odbiorcy to [...]	Opis niech przyjmie formę [...], w pierwszej [...]

Źródło: Opracowanie własne struktury prompta wg Macieja Chrzanowskiego (Wojewodzic, 2024, s.22)

Inny schemat prompta może być konstruowany w następujący sposób (Wojewodzic, 2024, s.24):

1. Jasne i konkretne polecenie: Powiedz dokładnie, co chcesz, aby AI zrobiło.
2. Podaj przykład: Pokaż, jak ma wyglądać odpowiedź, a potem poproś o podobną.
3. Ogranicz zakres odpowiedzi Powiedz, jak długą odpowiedź chcesz otrzymać lub co ma być w niej zawarte.
4. Zadawaj pytania krok po kroku Zadaj jedno pytanie, a potem, na podstawie odpowiedzi, zadawaj kolejne.
5. Ustal styl odpowiedzi Poproś o odpowiedź w konkretnym stylu (np. zabawnym, formalnym).

To proste struktury, które pozwalają na eksperymentowanie z własnymi wersjami promptów, co pozwala na ich modyfikację i rozwijanie w zależności od potrzeb. Powyższe punkty pozostawiają dużą elastyczność i nie oznaczają, że każdy prompt powinien zawierać wszystkie elementy. Można testować różne sposoby promptowania i cyklach poszczególnych sesji pogłębiać i poszerzać polecenia, oceniając efekty takiego podejścia.

Ciekawą budowę promptów znajdziemy także poniżej. Pierwszy z nich zwany schematem deklaracyjnych instrukcji, wygląda tak (Wang, Yizhong, i in., 2022):

- definicja – zwięzła, kompletna definicja zadania;
- pozytywny przykład - wkład przedstawiający przykład pozytywnego rezultatu, wraz z krótkim uzasadnieniem, przedstawiającym wskaźniki, na których nam zależy;
- negatywny przykład - wkład przedstawiający przykład negatywnego rezultatu, wraz z krótkim uzasadnieniem, przedstawiającym wskaźniki, na których nam zależy;
- następnie podaj przykładowe dane wejściowe (*input*) oraz listę możliwych, poprawnych odpowiedzi (*outputs*), które posłużą do oceny działania modelu.

Jest to ustandaryzowana, recenzowana w EMNLP rama opisu zadań, powszechnie używana w badaniach nad „instruction-tuning”, która bezpośrednio przekłada się na strukturę efektywnego prompta (definicja + pozytywne/negatywne przykłady + ograniczenia).

Kolejnym jest ReAct - „Thought → Action → Observation” (Yao, Shunyu, i in. 2023):

- Rola/Zasady (opcjonalnie: „Masz dostęp do narzędzi X. Przestrzegaj formatu”)
- Format pracy w pętli:
 - Myśl: (krótki krok rozumowania)
 - Działanie: (konkretne działanie/wywołanie narzędzia, np. „Uwzględnij[...]”)
 - Obserwacja: (co zwróciło narzędzie/środowisko)
 - Finałowa odpowiedź: (zwięzła odpowiedź)

6.5. Wsparcie twórców genAI

W procesie posługiwania się promptami i kolejnymi odsłonami modeli firmy będące twórcami tych narzędzi nie pozostawiają nas samym sobie. W poniższej tabeli można znaleźć wytyczne przygotowane przez Anthropic, Google, Microsoft i OpenAI. Istnieją także narzędzia optymalizacyjne, wspierające użytkowników w redagowaniu wartościowych i logicznie spójnych promptów.

Tabela 6.4. Wsparcie twórców genAI w zakresie tworzenia promptów

Anthropic	Google	Microsoft	OpenAI
Wskazówki			
Jasne, bezpośrednie polecenie wsparte przykładami. Stosowanie rozumowania krok po	Pisz naturalnym językiem, używając pełnych zdań. Nadaj rolę (personę), aby model przyjął	Rozpocznij od jasnego celu, który chcesz osiągnąć.	Jasne i precyzyjne polecenie: wyraźnie określ, co model ma zrobić.

kroku (Chain-of-Thought). Strukturyzacja promptu poprzez użycie znaczników (np. XML). Nadanie roli (persona), przygotowanie wstępnej odpowiedzi (<i>prefill</i>), a także umożliwienie modelowi przyznania się do braku wiedzy.	odpowiednią perspektywę. Podaj kontekst i jasno wskaż oczekiwaną formę odpowiedzi. Zachowuj zwięzłość i precyzję, a następnie stopniowo udoskonalaj prompt.	Dodaj kontekst, który pomoże ukierunkować model. Określ oczekiwania co do formy i treści odpowiedzi. W miarę potrzeby wskaż źródła lub dodatkowe ograniczenia.	Kontekst: dodaj dodatkowe informacje, które pomogą w lepszym zrozumieniu zadania. Styl: określ ton wypowiedzi (np. formalny, przyjazny). Iteracja: traktuj promptowanie jako proces powtarzalny – testuj i ulepszaj polecenia.
Adresy poradników			
https://docs.anthropic.com/en/docs/build-with-claude/prompt-engineering/overview	https://services.google.com/fh/files/misc/gemini-for-google-workspace-prompting-guide-101.pdf	https://support.microsoft.com/en-us/topic/learn-about-copilot-prompts-f6c3b467-f07c-4db1-ae54-ffac96184dd5	https://help.openai.com/en/articles/10032626-prompt-engineering-best-practices-for-chatgpt
Narzędzia optymalizacji promptów			
https://docs.anthropic.com/en/docs/build-with-claude/prompt-engineering/prompt-improver	https://cloud.google.com/blog/products/ai-machine-learning/announcing-vertex-ai-prompt-optimizer	Brak	https://platform.openai.com/chat/edit?models=gpt-5&optimize=true

Źródło: Opracowanie własne wg poradników Anthropic, Google, Microsoft, OpenAI

6.6. Budowanie promptów w pracy naukowej oraz ograniczanie halucynacji

Dla doktoranta prompt nie jest „pytaniem do czatu”, lecz specyfikacją zadania badawczego. Powinien jasno określać, co ma zostać wykonane, na jakich danych, w jakim formacie i według jakich kryteriów jakości. W większości zastosowań naukowych wystarczy czteroelementowy szkielet: opis danych wejściowych, precyzyjnie sformułowane zadanie, jednoznacznie zdefiniowana forma wyjścia oraz krótka procedura weryfikacji. Gdy praca jest bardziej złożona (na przykład konsolidacja wielu artykułów, przygotowanie rozdziału pracy), warto dołączyć kontekst i ograniczenia: długość, styl, zakazy (na przykład bez cytowań), a także metryki jakości, którymi ocenimy rezultat (kompletność, spójność, precyzja). Dobrą praktyką

jest dołączenie jednego krótkiego przykładu oczekiwanego wyjścia ułatwia on utrzymanie formatu w kolejnych iteracjach.

Kluczowe jest jedno: prompt ma zdejmować z modelu konieczność domysławiania się czegokolwiek istotnego. Jeśli przesyłasz pełny tekst artykułu, powiedz, które elementy są priorytetem (np. pytanie badawcze, próba, metody, wyniki główne), jaka ma być kolejność sekcji i czy dopuszczalne są uzupełnienia o wiedzę ogólną. Weryfikację należy wpisać wprost do treści i poprosić o listę braków lub niepewności, jeśli jakieś informacje nie występują w źródle. Taki bezpiecznik znacząco ogranicza ryzyko potencjalnych halucynacji.

Przykład 1. Szablon minimum do streszczenia rozdziału

```
## Dane wejściowe: pełny tekst rozdziału (poniżej między znacznikami).

## Zadanie: streść treść w 150–180 słowach dla czytelnika akademickiego.

## Forma wyjścia: pojedynczy akapit, język formalny, bez cytowań.

## Weryfikacja: na końcu dodaj 1–2 zdania „Braki/niepewności”, jeśli czegoś nie było w tekście.

<< {tu_wklej_tekst} >>
```

Źródło: opracowanie własne

Oprócz prostych streszczeń, można także sformułować zaawansowany prompt do generowania abstraktów, które są najkrótszym, ale najważniejszym gatunkiem tekstu naukowego. Dla doktoranta kluczowe jest panowanie nad strukturą, długością i tonem, tak by streszczenie było jednocześnie kompletne i zwarte. Poniżej przedstawiono gotowy do natychmiastowego użycia prompt do generowania abstraktów na podstawie treści artykułu.

Przykład 2. Szablon do generowania abstraktów

```
##Cel zadania

Na podstawie dostarczonych informacji wygeneruj abstrakt naukowy zgodny ze strukturą publikacji akademickich. Styl formalny, bez osobistych zaimków.

##Dane wejściowe

<< {tu_wklej_tekst_artykułu_albo_nazwę_zalaczonego_pliku} >>
```

##Struktura treści (merytoryka)

Uwzględnij kolejno:

1. Cel: luka badawcza i sposób jej wypełnienia.
2. Projekt/Metodyka: typ badania; źródła danych (np. Scopus, WoS); narzędzia (np. AI); standardy (np. PRISMA).
3. Wyniki/Wnioski: najważniejsze efekty i ich znaczenie, w stronie czynnej.
4. Ograniczenia: metodologiczne i merytoryczne, opisane rzeczowo.
5. Zastosowanie praktyczne: kto i jak może użyć wyników.
6. Oryginalność/Wartość poznawcza: co nowego wnosi praca.

##Format wyjścia

- Jeden spójny akapit o długości 250–300 słów (bez punktów i nagłówków).
- Język precyzyjny, rzeczowy; bez cytowań i odwołań bibliograficznych.

##Weryfikacja

Jeśli w danych brakuje istotnych informacji, po akapicie dodaj w jednym wierszu:
„Braki/niepewności: ...” (maks. 3 pozycje).

##Wynik

[Wstaw sam abstrakt (i ewentualną linię „Braki/niepewności”). Nie kopiuj nagłówków z promptu.]

Źródło: opracowanie własne

Modele językowe mają tendencję do mieszania wyjścia fragmentami instrukcji, zwłaszcza gdy te są zapisane luzem w tekście. Nagłówki ## działają jak sygnały segmentacji: jasno wskazują części polecenia (cel, wejście, zasady, format, wynik). Taki układ ma trzy praktyczne zalety. Po pierwsze, pozwala utrzymać porządek w długich poleceniach, łatwo dodać lub usunąć fragment bez naruszania reszty. Po drugie, ogranicza ryzyko, że model wstawi do odpowiedzi techniczne frazy („użyj strony czynnej”), bo są one odseparowane od sekcji „Wynik”. Po trzecie, sprzyja reprodukowalności: kiedy prompty przechowywane są w repozytorium, stała kolejność sekcji ułatwia wersjonowanie i testy A/B.

Aby separacja działała, warto przestrzegać czterech reguł. Po pierwsze, konsekwentna kolejność: Cel → Dane wejściowe → Zasady/Struktura → Format → Weryfikacja → Wynik.

Po drugie, jasne delimitery danych (<<...>>), dzięki którym model odróżnia treść źródłową

od instrukcji. Po trzecie, zakaz kopiowania nagłówków do odpowiedzi, należy wpisać to wprost w sekcji „Wynik”. Po czwarte, dodaj zawór bezpieczeństwa: należy dodać prośbę o wypunktowanie braków, kiedy danych nie ma. W praktyce to właśnie ta linia najbardziej ogranicza ryzyko halucynacji.

6.7. Dzielenie pracy na etapy w promptach

W praktyce naukowej najlepsze prompty nie delegują modelowi wykonania wszystkiego naraz. Zamiast tego rozbijają zadanie na czytelne etapy i prowadzą model przez kolejne kamienie milowe: od ekstrakcji faktów, przez syntezę i szkic, aż po weryfikację. Taki układ zwiększa trafność, ułatwia kontrolę jakości i pozwala zatrzymać się po każdym kroku, aby coś doprecyzować. Poniżej opisana została zwięzła metoda etapowania, zasady zapisu poleceń z nagłówkami **##** oraz kompletne, gotowe do użycia przykłady w tym pełny pipeline do tworzenia abstraktu naukowego.

Etapowanie rozwiązuje trzy typowe problemy: przeciążenie modelu, brak powtarzalności, limity tokenów w output oraz halucynacje. Gdy model dostaje małe, precyzyjne kroki (na przykład najpierw wyodrębnij metody, dopiero potem generuj abstrakt), ryzyko domysłów spada, a wynik jest łatwiejszy do zweryfikowania. Dodatkowo po każdym etapie można podjąć decyzję: „kontynuuj”, „popraw”, „zmień kryteria”. W pracy doktoranta oznacza to realną oszczędność czasu przy zachowaniu kontroli merytorycznej.

Instrukcje należy dzielić na bloki z nagłówkami **##**. Taki zapis klarownie oddziela cel, wejście, kroki, format wyjścia i weryfikację, a do każdego kroku można dodać doprecyzowania. Dane źródłowe można zamykać w delimiterach `<<<...>>>`, by model nie mieszał ich z instrukcjami. W blokach „Wynik” warto dopisać, by nie kopiować nagłówków do odpowiedzi oraz by zatrzymać się i zapytać, czy kontynuować.

Przykład 3. Prompt dzielący pracę na etapy

##Cel

Krótko: co jest celem i dlaczego (1–2 zdania).

##Dane wejściowe

`<< {tu_wklej_tekst/tabelę} >>`

##Etap 1 - Ekstrakcja

Instrukcja: co wydobyć z wejścia i w jakim formacie (na przykład plik csv/tabela).

Wynik: [tu opis formatu].

Po zakończeniu zapytaj: „Czy kontynuować do Etapu 2? [T/N]”. Zatrzymaj się.

##Etap 2 – Synteza/Plan

Instrukcja: ułóż konspekt/mapę treści z elementów z Etapu 1.

Wynik: [np. lista punktów].

Zapytaj o zgodę na Etap 3. Zatrzymaj się.

##Etap 3 – Szkic/Wersja robocza

Instrukcja: napisz pierwszą wersję, trzymając się konspektu. Wynik: [np. 250–300 słów, jeden akapit].

Zapytaj o zgodę na Etap 4. Zatrzymaj się.

##Etap 4 – Weryfikacja i poprawa

Instrukcja: wypisz braki/ryzyka (max 3) i wygeneruj poprawioną wersję.

Wynik: [finalny tekst + „Braki/niepewności”, jeśli występują].

Źródło: opracowanie własne

Stworzenie spójnego i zgodnego ze standardami abstraktu wymaga nie tylko znajomości treści artykułu, ale również umiejętności pracy krok po kroku. Dlatego przygotowany został poniższy schemat na bazie schematu wyżej, ale rozszerzony i dostosowany o specyficzne elementy dotyczące abstraktu. Podział pracy na etapy pozwala oddzielić wydobywanie danych od ich syntezy i redakcji, dzięki czemu ryzyko pominięcia lub halucynacji jest minimalne.

Przykład 4. Rozszerzony schemat do pracy nad abstraktem

Cel

Na podstawie danych o badaniu przygotuj finalny abstrakt zgodny z praktyką akademicką.

Dane wejściowe

<< {tu_wklej notatki albo skrót artykułu} >>

##Etap 1 – Ekstrakcja pól

Z wejścia wyodrębnij dokładnie i bez dopowiadania:

- cel badania (jaka luka i jak ją wypełniamy),
- projekt/metodyka (typ badania; źródła danych: np. Scopus/WoS; narzędzia: np. AI; standardy: np. PRISMA),

- wyniki kluczowe (najważniejsze liczby/efekty),
- wnioski (znaczenie wyników),
- ograniczenia (metodologiczne i merytoryczne),
- zastosowania praktyczne,
- oryginalność/wartość poznawcza.

Wynik zwróć jako obiekt JSON z kluczami: cel, metodyka, wyniki, wnioski, ograniczenia, zastosowania, oryginalność.

Po zwróceniu JSON napisz: „Etap 1 zakończony. Czy kontynuować do Etapu 2? [T/N]”.

Zatrzymaj się.

##Etap 2 – Konspekt abstraktu

Na podstawie JSON ułóż 6-punktowy plan treści w kolejności:

1) Cel, 2) Metodyka, 3) Wyniki, 4) Wnioski, 5) Ograniczenia, 6) Zastosowania i oryginalność.

Każdy punkt maks. 20 słów. Nie wprowadzaj nowych informacji.

Na końcu zapytaj: „Etap 2 zakończony. Kontynuować do Etapu 3? [T/N]”. Zatrzymaj się.

##Etap 3 – Szkic abstraktu

Napisz jeden spójny akapit 250–300 słów w stylu formalnym, bez nagłówków i bez cytowań.

Trzymaj się kolejności z konspektu; używaj strony czynnej i precyzyjnych sformułowań.

Na końcu zapytaj: „Etap 3 zakończony. Uruchomić weryfikację (Etap 4)? [T/N]”. Zatrzymaj się.

##Etap 4 – Weryfikacja i korekta

1. Wypisz „Braki/niepewności” (max 3), jeśli jakiejś informacji brakuje w danych wejściowych.

2. Zredaguj wersję poprawioną (ten sam limit słów), korygując usterki stylistyczne i logiczne.

Zwróć:

- „Braki/niepewności: ...” (albo „brak”)

- „Abstrakt – wersja finalna: ...”

Źródło: opracowanie własne

Praca nad abstraktem w czterech etapach (ekstrakcja, konspekt, szkic i weryfikacja) pozwala stopniowo budować tekst, jednocześnie zachowując kontrolę nad jego zawartością i formą. Rozdzielenie zadań na mniejsze kroki zmniejsza ryzyko błędów, a sekcja końcowej weryfikacji pełni rolę zaworu bezpieczeństwa, wskazując potencjalne braki i nadając tekstowi

ostateczną spójność. Dzięki temu uzyskać można narzędzie, które nie tylko przyspiesza pracę, ale też uczy systematycznego i metodycznego podejścia do pisania naukowego.

Zadanie do wykonania

Należy wziąć jeden artykuł naukowy (na przykład z własnej bibliografii doktorskiej) i zastosuj powyższy pipeline w praktyce:

Etap 1 – Ekstrakcja: wyodrębnij z artykułu podstawowe informacje i zapisz je w formacie JSON.

Etap 2 – Konspekt: na podstawie JSON przygotuj 6-punktowy plan abstraktu (po maks. 20 słów na punkt).

Etap 3 – Szkic: napisz jeden akapit (250–300 słów) w stylu formalnym, zgodny z kolejnością z konspektu.

Etap 4 – Weryfikacja: wskaż maks. trzy „Braki/niepewności”, a następnie popraw szkic, tworząc wersję finalną abstraktu.

Cel zadania: porównać szkic i wersję finalną, aby zobaczyć, jak etap weryfikacji wpływa na jakość i precyzję abstraktu.

Podsumowanie rozdziału 6

Konstruowanie własnego stylu promptowania, dostosowanego do indywidualnych potrzeb jest ustawicznym procesem. Swoje umiejętności można rozwijać również w dialogu z chatbotami prosząc o wskazówki i otrzymując informacje zwrotne dotyczące struktur używanych promptów. Dobre rozwiązania i prompty, który przynoszą zadawalające efekty warto zapisywać w osobistej bibliotece, aby móc do nich sięgać i rozwijać w dalszej pracy. Rozwój kompetencji w zakresie inżynierii promptów wiąże się bezpośrednio z rozwojem warsztatu pracy dydaktycznej i naukowej młodego naukowca. Doświadczenie i wyciąganie wniosków z praktyki integrowania genAI jest ustawicznym cyklem refleksji i modyfikacji metod i narzędzi w stosowanych w ramach procesu promptowania.

7. AI w pisaniu publikacji naukowych

Sztuczna inteligencja coraz częściej wspiera badaczy w procesie pisania publikacji naukowych od przeglądu literatury, przez redakcję językową, po automatyzację wybranych etapów analizy. Narzędzia generatywne, takie jak duże modele językowe, oferują możliwość szybszego opracowywania szkiców, porządkowania struktury tekstu czy tworzenia streszczeń.

Rozdział ten przedstawia przykłady zastosowań AI w pisaniu i redagowaniu prac naukowych, omawia różne strategie współpracy z tymi systemami oraz wskazuje dobre praktyki zapewniające rzetelność i przejrzystość publikacji.

7.1. Narzędzia AI wykorzystywane w pracy naukowej

ChatGPT i inne chatboty oparte na LLM

Na temat narzędzi takich jak ChatGPT, Gemini, Claude, czy Copilot w tej publikacji można znaleźć bardzo wiele informacji. W związku z ich ustawicznym rozwojem w tym miejscu podkreślone zostaną te funkcjonalności, które mogą wspierać uzyskiwanie wartościowych rezultatów. Skupimy się na ChatGPT ponieważ posiada on najwięcej rozwiązań, które na ten moment nie są jeszcze dostępne w pozostałych narzędziach. ChatGPT posiada w menu użytkownika sekcję pozwalającą na udzielenie czatowi informacji, które umożliwiają na spersonalizowanie użytkownika. Specjalny formularz pozwala na wpisanie naszego doświadczenia zawodowego, specyficznych życzeń związanych z wymaganiami dotyczącymi tego, jak chat powinien się z nami komunikować, jaki styl komunikacji nas interesuje. Uzupełnienie tych informacji pozwala głębiej skalibrować chat/model do naszych potrzeb.

Dodatkową istotną funkcjonalnością jest pamięć. ChatGPT może zapamiętywać i gromadzić informacje o nas w taki sposób, aby mieć do nich dostęp w czasie konwersacji z użytkownikiem. Zapisane dane możemy znaleźć w menu użytkownika, w funkcji personalizacja, w funkcji *zarządzanie pamięcią*. Dodatkowo w sekcji tej można włączyć lub wyłączyć funkcję, która pozwala modelowi na dostęp do danych znajdujących się w różnych naszych konwersacjach i na ich podstawie generuje odpowiedzi. Dlatego, aby dbać o wartość odpowiedzi, można korzystać też z konwersacji tymczasowych, które nie zapisują się w sekcji rozmów, dzięki czemu w przypadku rozmów, na których nie chcemy zanieczyszczać zbioru udostępnianych przez nas danych, można w ten sposób realizować rozmowy z chatem.

Ponieważ narzędzia te rozwijają się bardzo intensywnie, w zakresie funkcjonalności, które są dostępne dla użytkowników, możliwe, że w innych usługach pojawiają się siostrzane rozwiązania, a ich stosowanie jest rekomendowane w celu otrzymywania lepszych rezultatów wynikających z personalizacji i kalibracji modeli.

NotebookLM

NotebookLM to jedno z narzędzi Google, które jest wspierane przez model Gemini. Różni się tym od typowych rozwiązań, że pozwala na pracę z własnymi źródłami danych i swoje odpowiedzi ogranicza tylko do tych zasobów, które zostaną udostępnione przez użytkownika.

W wersji bezpłatnej użytkownik może pracować jednocześnie z 50 dokumentami, każdy z nich ma limit w liczbie 500 000 znaków. NotebookLM obsługuje różne formaty plików, w tym:

- pliki tekstowe: PDF, .txt, Markdown, Google Docs, Google Slides;
- pliki audio: MP3, WAV, M4A, MP4, AIFF, AAC, CAF, AMR (plik audio jest transkrybowany, a tekst zapisywany jako źródło);
- treści internetowe: linki do stron internetowych (importowany jest tylko tekst) oraz linki do publicznych filmów na YouTube (importowana jest transkrypcja);
- inne: skopiowany i wklejony tekst.

Generowane przez Notebooka odpowiedzi nie tylko bazują na udostępnionych materiałach, ale również wskazują poprzez interaktywne odnośniki te fragmenty zasobów, na podstawie których generowane są odpowiedzi. Pozwala to zminimalizować ryzyko halucynacji, a użytkownikom daje większą kontrolę nad uzyskiwanymi rezultatami.

Dodatkowo NotebookLM pozwala między innymi na generowanie podcastów i videocastów ilustrowanych infografikami. Zasoby te stanowią syntezę udostępnionych zasobów, a treści mogą być realizowane również w języku polskim. Pozwala to więc na uzyskanie materiałów, które są nie tylko tekstowe, ale pozwalają na zapoznanie lub utrwalenie określonego zakresu wiedzy z wykorzystaniem różnych modalności.

Scite.ai - Smart Citations i asystent AI w pracy młodego naukowca

Scite.ai stanowi platformę badawczą opartą na sztucznej inteligencji, której kluczowym wyróżnikiem jest baza Smart Citations. Każde cytowanie w tej bazie jest klasyfikowane jako wspierające, kontrastujące lub jedynie wzmiankujące określony wynik, co pozwala umieścić pracę w kontekście całego korpusu literatury. Rozwiązanie łączy dostęp do treści otwartych i płatnych z przetwarzaniem pełnych tekstów, a moduł Assistant by scite generuje odpowiedzi z jawnie widocznymi odwołaniami do źródeł, ograniczając ryzyko halucynacji modeli językowych. Z narzędzia korzystają indywidualni badacze i instytucje (uczelnie, wydawnictwa, firmy), a sam produkt rozwijany jest od 2018 r. i obecnie część Research Solutions oferuje również przeglądarkowe rozszerzenie, pulpity analityczne, sprawdzanie bibliografii oraz wyszukiwarkę zdań cytujących.

Wykorzystanie scite w badaniach literaturowych przebiega w trybie konwersacyjnym: użytkownik formułuje zapytanie w języku naturalnym, a Asystent odpowiada syntetycznie i równocześnie dołącza listę publikacji wraz z DOI i kontekstem cytowania. Dzięki temu już na poziomie pojedynczego akapitu można ocenić, czy powołania są zgodne z wnioskami, czy też stanowią jedynie neutralne wzmianki. Funkcjonalność Citation Statement Search umożliwia

dodatkowo wyszukiwanie bezpośrednio po zdaniach cytujących zaczerpniętych z pełnych tekstów, co ułatwia szybkie identyfikowanie konkretnych dowodów, na przykład opisów metod, wyników lub ograniczeń. W ramach jednego środowiska można tworzyć kolekcje artykułów, nadawać im etykiety, uzyskiwać zagregowane statystyki oraz włączać powiadomienia o nowych cytowaniach i publikacjach z danego obszaru.

Skonfigurowanie sposobu pracy Asystenta odgrywa zasadniczą rolę dla jakości wyników. Rysunek 1 przedstawia panel Assistant Settings, w którym można określić wymóg dołączania referencji (zawsze zalecane), rodzaj źródeł dowodów (tylko abstrakty, wyłącznie zdania cytujące lub oba strumienie łącznie), sekcje artykułu brane pod uwagę jako dowody (na przykład Methods lub Results), tryb tabelaryczny odpowiedzi, przedział lat, typy publikacji (recenzowane artykuły, przeglądy, preprinty), styl cytowania (na przykład APA), model generatywny, długość odpowiedzi, liczbę publikacji konsultowanych w odpowiedzi oraz kryterium rankingu referencji.

The screenshot shows the 'Assistant Settings' interface. It features several sections with radio buttons, dropdown menus, and sliders. The 'Specify Reference Requirement' section has three options: 'Let Assistant decide' (selected), 'Always use references', and 'Never use references'. The 'Specify Evidence Source' section has three options: 'Both' (selected), 'Abstracts only', and 'Citation Statements only'. The 'Specify Evidence Sections' section has a dropdown menu labeled 'Select...'. The 'Use Table Mode' section has a toggle switch. The 'Reference Year Range' section has 'From' and 'To' input fields. The 'Publication Types' section has a dropdown menu labeled 'Select...'. The 'Citation Style' section has a dropdown menu labeled 'APA'. The 'Model' section has a dropdown menu labeled 'GPT4 o-mini'. The 'Response Length' section has a dropdown menu labeled 'Medium'. The '# Publications to consult' section has a slider set to 40. The 'Reference Ranking' section has a dropdown menu labeled 'Relevance'. At the bottom, there are 'Clear All' and 'Apply' buttons.

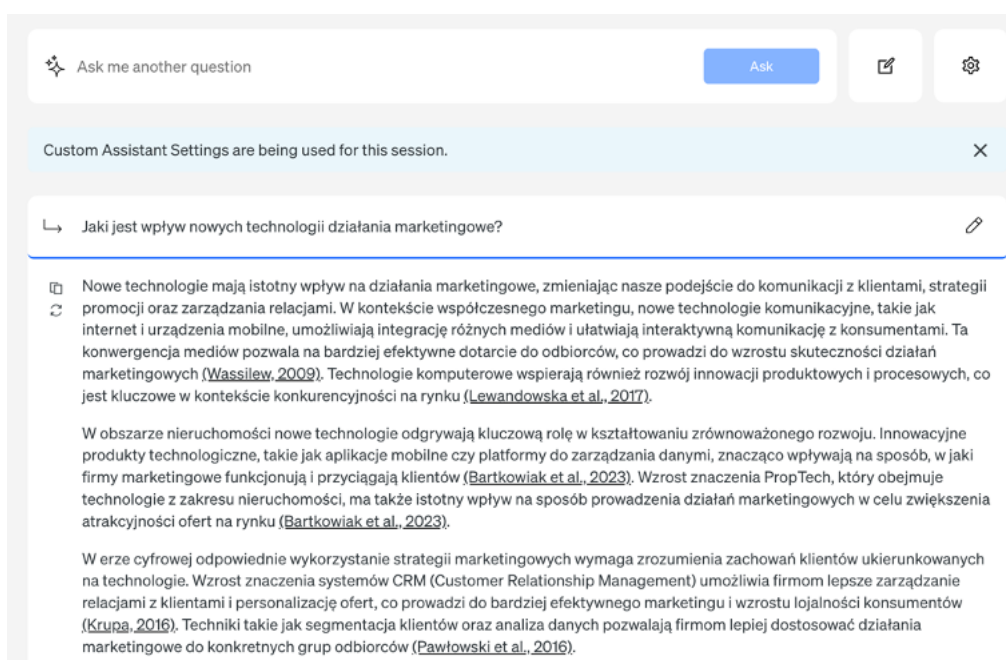
Rys. 7.1. Ustawienia Asystenta scite.ai panel „Assistant Settings”, z najważniejszymi przełącznikami.

Źródło: scite.ai

Wprowadzenie w ustawieniach zaprezentowanych na rysunku 1 ograniczeń rocznikowych oraz typologicznych pozwala precyzyjnie sterować aktualnością i

wiarygodnością wyników, natomiast zwiększenie parametru Publications to consult zwykle podnosi kompletność syntezy kosztem czasu generowania. Dodatkowe pole Journals umożliwia ograniczenie kwerendy do wybranych czasopism, co bywa przydatne przy tworzeniu mapy literatury dla wąskich dziedzin.

Odpowiedź Asystenta ma postać spójnych akapitów z wbudowanymi odnośnikami do literatury i można ją traktować jako zwięzły przegląd stanu badań. Rysunek 7.2 ilustruje tę formę: widoczne są akapity streszczające najważniejsze wątki oraz dyskretne odwołania bibliograficzne, które po rozwinięciu prowadzą do pełnego kontekstu cytowania w Smart Citations. Dalsze przejrzanie kilku kluczowych pozycji pozwala szybko ocenić, które tezy są szeroko wspierane, a które podlegają sporom. Informacja o charakterze cytowania (wspierające, kontrastujące, wzmiankujące) ułatwia to priorytetyzację pozycji literaturowych i decyzje dotyczące cytowania. W ustawieniach odpowiedzi możliwe jest włączenie trybu tabelarycznego dla porównań metod, populacji czy metryk, jednak nawet w trybie tekstowym Asystent dąży do klarownej struktury tematycznej.



Rys. 7.2. Przykładowy wynik odpowiedzi Asystenta scite.ai na pytanie badawcze

Źródło: scite.ai

Transparentność doboru źródeł została rozwiązana poprzez odrębną kartę z referencjami oraz strategią wyszukiwania. Rysunek 7.3 prezentuje numerowaną bibliografię z identyfikatorami DOI oraz szczegóły procedury kwerendy zastosowanej przez Asystenta (między innymi użyte frazy i filtry). Taka dokumentacja ułatwia replikację procesu wyszukiwania w pracach przeglądowych i metodycznych, a także usprawnia komunikację z promotorem i recenzentami.

References	Search Strategy
1. (2009). <u>Information and communication technology in contemporary society</u>. <i>gospodarka narodowa</i>, 236(11-12), 67-95. https://doi.org/10.33119/gn/101241 2. (2023). <u>Nowe technologie na rynku nieruchomości – w poszukiwaniu zrównoważonego rozwoju</u>. https://doi.org/10.18559/978-83-8211-156-9 3. (2017). <u>Innovation complementarity and new-product exports</u>. <i>gospodarka narodowa</i>, 287(1), 95-117. https://doi.org/10.33119/gn/100747 4. (2025). <u>Klient i komunikacja marketingowa w erze cyfrowej, ujęcie praktyczne...</u> https://doi.org/10.60154/klientikomunikacja/2025 5. (2015). <u>Relacje między strategią marketingową a strategią finansowania</u>. <i>zeszyty naukowe sggw polityki europejskie finanse i marketing</i>, 13(62), 30-39. https://doi.org/10.22630/pefim.2015.13.62.3 6. (2016). <u>The marketing communication of university as a chance of gaining competitive advantage</u>. <i>marketing i zarządzanie</i>, 45, 193-201. https://doi.org/10.18276/miz.2016.45-17	

Rys. 7.3. Lista źródeł i karta Search Strategy

Źródło: scite.ai

Warto podkreślić, że scite umożliwia dołączanie własnych materiałów w zakładce Sources. Po ich dodaniu Asystent włącza te dokumenty do puli publikacji i może je jawnie cytować obok zasobów zewnętrznych, co sprzyja integrowaniu wyników przeglądu literatury z raportami projektowymi, preprintami czy publikacjami autora.

Z punktu widzenia warsztatu doktoranta narzędzie wspiera kilka typowych scenariuszy. W szybkim rekonesansie tematu rekomendowane jest ustawienie obowiązkowych referencji i uwzględnienie zarówno abstraktów, jak i zdań cytujących, przy równoczesnym ograniczeniu lat publikacji, co zapewnia aktualność. W przygotowaniu przeglądu zakresu literatury pomocne bywa zawężenie typów do prac recenzowanych i przeglądowych oraz użycie trybu tabelarycznego, który ułatwia porównania. Na etapie finalizacji rękopisu funkcja Reference Check pozwala ocenić jakość i aktualność cytowań oraz szybko wykryć powołania na wyniki kontrowersyjne. Przy pytaniach metodologicznych zaleca się wskazanie w ustawieniach sekcji Methods oraz wymuszenie informacji wyłącznie ze zdań cytujących, co zwiększa precyzję odniesień.

Korzystanie z Smart Citations sprzyja krytycznej lekturze, lecz nie zwalnia z obowiązku weryfikacji kontekstu. Każde powołanie powinno zostać sprawdzone bezpośrednio w tekście źródła, zwłaszcza gdy klasyfikowane jest jako kontrastujące. W praktyce badawczej warto także

dopasować końcową bibliografię do wymogów uczelni lub czasopisma, mimo że scite oferuje formatowanie w popularnych stylach, takich jak APA.

Całość infrastruktury została zaprojektowana z myślą o obserwowalności procesu generowania odpowiedzi i zgodności z prawem autorskim, co przejawia się w cytowaniu do pełnych tekstów oraz ekspozycji zdań cytujących zamiast kopiowania obszernych fragmentów.

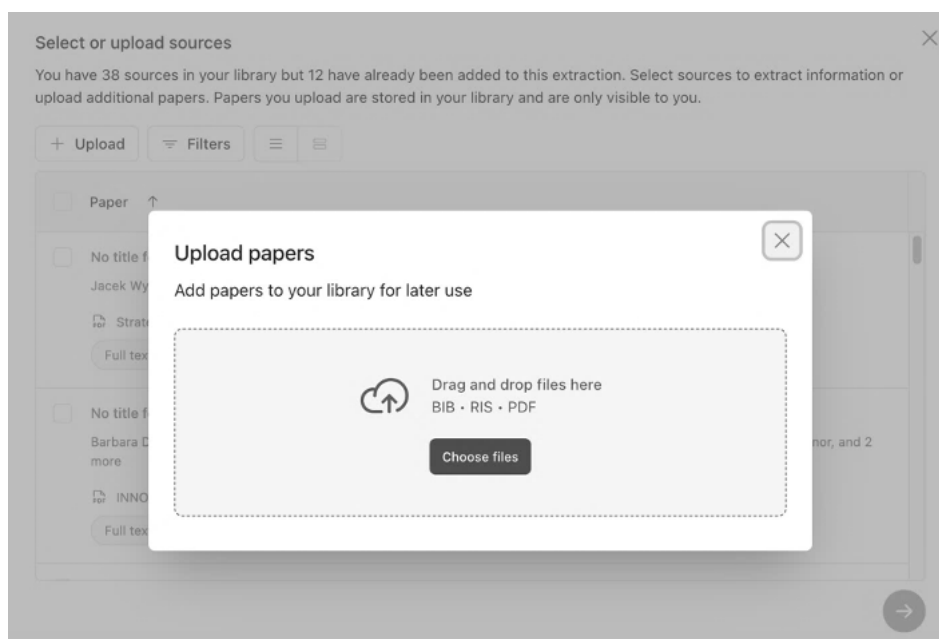
Scite integruje wyszukiwanie pełnotekstowe, klasyfikację cytowań oraz asystenta konwersacyjnego w jednym, weryfikowalnym środowisku pracy. Pozwala to szybciej mapować literaturę, świadomie zarządzać bibliografią i łączyć wyniki przeglądu z własnymi materiałami.

Elicit - asystent przeglądów literatury i ekstrakcji danych

Elicit jest narzędziem wspierającym badaczy w automatyzacji zadań wymagających pracy z dużymi zbiorami publikacji naukowych. Oprogramowanie przyspiesza streszczanie artykułów, ekstrakcję danych, selekcję literatury oraz syntetyzowanie wniosków. Zgodnie z deklaracjami producenta przeglądy systematyczne realizowane w Elicit mogą zajmować nawet o 80% mniej czasu, a każda odpowiedź jest podparta cytatem z pracy źródłowej.

System wyszukuje w bazie ponad 125 milionów publikacji i umożliwia natychmiastowe dodanie do 500 pokrewnych artykułów do własnego zbioru. Szczególnie istotną cechą jest przejrzystość ścieżki dowodowej: od odpowiedzi można przejść do fragmentu tekstu w PDF, co ułatwia weryfikację i audyt procesu, który wykonuje niniejsze oprogramowanie.

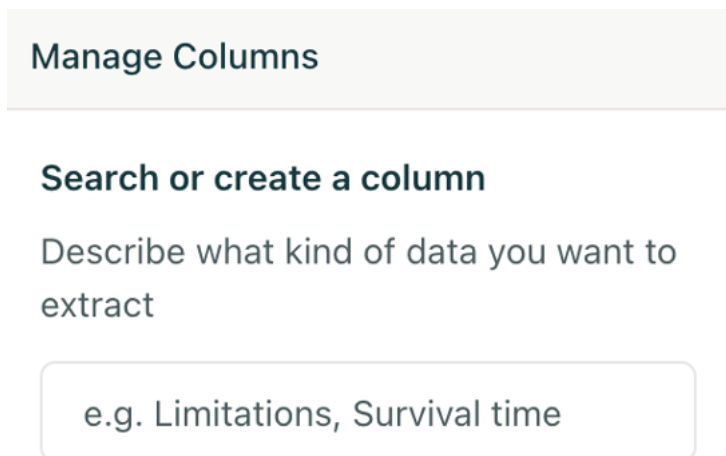
Praca rozpoczyna się od zbudowania prywatnej biblioteki źródeł. Artykuły dodaje się przez przeciągnięcie plików PDF, BIB lub RIS lub wybór z dysku, okno odpowiedzialne za import widoczne jest na Rysunku 7.4.



Rys. 7.4. Upload papers - dodawanie PDF/BIB/RIS do prywatnej biblioteki Elicit.

Źródło: elicit.com

Po załadowaniu materiału definiuje się, jakie informacje mają zostać wyodrębnione z prac. Służy do tego panel Manage Columns, co przedstawiono na rysunku 7.5, w którym można wyszukać istniejące kolumny (na przykład Limitations, Outcome measured) albo utworzyć własne.



Rys. 7.5. Manage Columns - tworzenie lub wybór kolumn do ekstrakcji danych.

Źródło: elicit.com

Każda kolumna posiada pole Instructions, czyli miejsce na precyzyjny prompt opisujący sposób ekstrakcji. Na rysunku 7.6 przedstawiono formularz edycji takiej kolumny wraz z przykładową instrukcją wymuszającą krótkie, jedno sentencyjne punkty podsumowujące wyniki badań.

 Edit custom column

Tell Elicit more about this column to improve accuracy. [Get guidance and examples here.](#)

Column name

Main findings

Instructions (optional)

Summarize the results or conclusions of the study. Use bullet points (each starting with a dash). Each bullet point should consist of only one concise sentence. Give a minimum of one bullet point and a maximum of three. Make sure they convey the most important takeaways from the paper. Avoid being redundant.

Answer Structure

Any answer Specified Yes/No/Maybe

Cancel Save

Rys. 7.6. Edit custom column - pole Instructions do formułowania własnych promptów

Źródło: elicit.com

Praktykę iteracji promptów ułatwia menu kolumn (edytowanie, duplikowanie, filtrowanie i zapisywanie jako preset), widoczne na rysunku 7.7.



of c
are
ist
is c
g A
tic
nti
marketing: algorithmic, societal, and

 Edit

 Duplicate

 Filter

 Save as preset

 Delete column

Instructions (optional)

Summarize the results of the study. Use bullet points (each starting with a dash). Each bullet point should consist of only one concise sentence. Give a minimum of one bullet point and a maximum of three. Make sure they convey the most important takeaways from the paper. Avoid being redundant.

Answer Structure

Any answer Spe

Rys. 7.7 Menu kolumny - edycja, duplikowanie, filtrowanie, zapisywanie jako preset

Źródło: elicit.com

Format odpowiedzi można dodatkowo ustandaryzować, korzystając z opcji Answer Structure. Wybór Any answer pozostawia modelowi swobodę w ramach instrukcji, opcja Specified pozwala określić ściśle strukturę, natomiast Yes/No/Maybe wspiera kodowanie

binarne lub wielokrotnego wyboru przy ocenie ryzyka biasu. Panel wyboru struktury odpowiedzi pokazano na rysunku 7.8.

Paper	Main findings	Methodology
☐ ETHICAL IMPLICATIONS OF ARTIFICIAL INTELLIGENCE MARKETING Nandyalu Lokeswar +3 INTERNATIONAL JOURNAL OF SCIENTIFIC RESEARCH IN ENGINEERING AND MANAGEMENT ETHICAL-IMPLICATIONS-OF-ARTIFICIAL-INTELLIGENCE 2024 · 13 citations	- Perceptions of data protection, fairness, transparency, and autonomy are interconnected in AI-driven marketing. - Consumer trust in autonomy is positively correlated with perceptions of data protection, fairness, and transparency. - The marketing AI tool has a significant impact on awareness, particularly among urban and younger demographics.	- Mixed-method study approach - Purposive sampling for diverse stakeholder representation - Collaboration with market research platforms for sampling - Sample size determined using Slovic's - Google Forms for online surveys
☐ Ethical Marketing AI? A Structured Literature Review Consumer Behavior Yiran Su +3 Ethical Marketing AI? A Structured Literature Review of Consumer Behavior Citations unknown	- The study identifies three clusters of ethical issues related to AI in marketing: algorithmic, societal, and existential. - The field of ethical marketing AI is still in its infancy, with ethical issues not yet fully addressed in the marketing domain. - There is a need for clear programs to mitigate the risks caused by these ethical dilemmas in marketing.	- Conducted a structured literature review related to AI in marketing and consumer behavior - Followed the literature review process and and Watson (2019). - Identified keywords related to AI ethics - Created a code book as a review protocol - Searched literature using keywords at the sample based on relevance. - Applied inclusion criteria to filter articles - Conducted a qualitative synthesis to synthesize selected papers.
☐ Ethical Considerations in AI-Based Marketing: Balance Trust Swati Sharma +6 Tuijin Jishu/Journal of Propulsion Technology	- Transparency and accountability are foundational pillars of ethical AI-based marketing. - Data privacy and informed consent are critical ethical concerns in AI-based marketing. - Ethical considerations are complementary to	- Not mentioned (the paper is a theoretical ethical considerations in AI-based marketing not describe a specific methodology)

Rys. 7.8. Tabela ekstrakcji - przykładowe wypełnienie kolumn Main findings i Methodology
Źródło: elicit.com

Wyniki ekstrakcji są prezentowane w tabeli, w której każdy wiersz odpowiada jednej publikacji, a kolumny wypełniane są treściami dostosowanymi do użytych instrukcji. Przykładowy widok takiej tabeli, z kolumnami Main findings i Methodology, przedstawiono na rysunku 7.8. Ten widok stanowi wygodne środowisko do dalszego sortowania, filtrowania i rozszerzania projektu o kolejne kolumny, w planach płatnych możliwy jest także eksport do CSV.

Elicit oferuje zestaw funkcji zaprojektowanych z myślą o pracy zgodnej z rygorami dowodowymi. Moduł przeglądów systematycznych prowadzi przez etapy wyszukiwania, ekranowania, ekstrakcji i syntezy. Filtrowanie może być częściowo zautomatyzowane: system generuje kryteria wraz z wstępnymi ocenami, które następnie podlegają weryfikacji. Ekstrakcja skaluje się do setek prac w krótkim czasie, w tym do informacji umieszczonych w tabelach PDF, co zwykle bywa przeszkodą przy ręcznej analizie. Dostępny jest także komponent raportowy Elicit Reports, który tworzy szybkie, edytowalne raporty oparte na zautomatyzowanym przeglądzie, z pełną możliwością przeglądu i korekty każdego kroku. Komunikacja z kilkoma plikami jednocześnie umożliwia zadawanie pytań przekrojowych i prowadzenie analizy porównawczej bez opuszczania interfejsu.

Efektywność Elicit zależy od jakości instrukcji definiowanych w polu Instructions. Dobrze zaprojektowany prompt powinien określać zakres (na przykład wyłącznie sekcje

Methods i Results), format (na przykład krótkie zdania w punktach lub z góry zdefiniowane pola), ograniczenia długości i oczekiwane jednostki miary, a także wymagać wskazania cytatu potwierdzającego odpowiedź. Zastosowanie takich zasad pozwala uzyskać porównywalne wiersze w tabeli oraz ogranicza rozbieżności między publikacjami. W praktyce użyteczne jest duplikowanie działających kolumn i ich iteracyjna modyfikacja, a następnie zapisywanie jako presetów dla kolejnych projektów. Im większą wiedzę w zakresie prompt engineeringu dysponuje użytkownik tego oprogramowania, tym lepsze rezultaty otrzyma. Stąd też do właściwego i profesjonalnego wykorzystania narzędzi do pracy naukowej, niezbędna jest podstawowa wiedza na temat bezpośredniej pracy z modelami językowymi, promptowaniem, co zostało opisane w rozdziale 1 części praktycznej niniejszej monografii.

Proces przygotowania przeglądu systematycznego przebiega w Elicit w sposób uporządkowany. Po zdefiniowaniu pytania badawczego należy zaimportować zgromadzone pliki PDF i rozszerzyć bazę o publikacje wyszukane w systemie. Następnie uruchamia się filtrowanie z wykorzystaniem wygenerowanych kryteriów, po czym planuje się zestaw kolumn odpowiadających na pytanie badawcze. Po pierwszej ekstrakcji wskazane jest sprawdzenie kilku wierszy wraz z cytatami i doprecyzowanie instrukcji. W dalszej kolejności warto skorzystać z opcji Specified lub Yes/No/Maybe, aby wymusić spójny format, a na koniec wyeksportować tabelę i wykorzystać ją do narracyjnej syntezy lub metaanalizy. Tak prowadzony proces jest transparentny i łatwy do audytu, ponieważ każdą informację można prześledzić do odpowiedniego fragmentu źródła.

Chociaż Elicit redukuje czasochłonność pracy, narzędzie opiera się na modelach językowych, które zwłaszcza przy niejednoznacznych plikach PDF lub słabym OCR mogą wymagać ręcznej weryfikacji. Kluczowe jest konsekwentne sprawdzanie cytatów i dopasowanie promptów do specyfiki dziedziny, w tym standaryzacja terminologii oraz jednostek. Lepsze wyniki zapewnia dzielenie złożonych zadań na kilka prostszych kolumn zamiast jednego, nadmiernie ogólnego polecenia. Zachowanie tych zasad pozwala w pełni wykorzystać potencjał narzędzia przy jednoczesnym utrzymaniu wysokiej jakości informacji, mających odzwierciedlenie w cytowanych pozycjach literaturowych.

Elicit stanowi profesjonalną platformę wspierającą badania oparte na literaturze. Integracja wyszukiwania, filtrowania, ekstrakcji i raportowania w jednym środowisku, a także możliwość ścisłego kontrolowania formatu odpowiedzi poprzez własne prompty i strukturę odpowiedzi, przekładają się na znaczącą oszczędność czasu przy zachowaniu przejrzystości metodologicznej.

7.2. Różne strategie współpracy z AI

Współpraca z systemami sztucznej inteligencji może przyjmować różne formy od twórczego treningu, przez wsparcie redakcyjne, po automatyzację wybranych etapów pracy badawczej. Przegląd różnych praktyk pokazuje, że najbardziej efektywne są modele, w których człowiek i maszyna uzupełniają się kompetencyjnie, a rola badacza przesuwa się z wykonywania żmudnych czynności w stronę kontrolowania procesu, kontroli jakości i formułowania wniosków (Tomczyk, Brüggemann i Vrontis, 2024).

Pierwsza strategia to traktowanie AI jako asystenta pisarskiego i redakcyjnego. Narzędzia generatywne pomagają porządkować strukturę tekstu, proponować ujęcia argumentów, a także wygładzać styl i dopasowywać go do konwencji dyscypliny. Kluczowa pozostaje jednak rola autora, to on nadaje kierunek, weryfikuje treści i dba o rzetelność, AI nie jest autorem publikacji i nie może go zastąpić (Haber, Przeglasińska i Jemielniak, 2025).

Druga strategia to użycie AI jako partnera w przeglądzie literatury. Badania porównawcze wskazują, że tradycyjne bazy, takie jak Scopus i Web of Science, wciąż wygrywają w precyzji i jakości wyników, natomiast narzędzia oparte na AI częściej odkrywają publikacje unikalne, które umykają klasycznym przeglądom. W praktyce najbardziej efektywne są podejścia łączone, w których wyniki z baz tradycyjnych są uzupełniane o sugestie systemów AI (Tomczyk, Brüggemann, Mergner i Petrescu, 2024; Tomczyk, Brüggemann, Mergner i Petrescu, 2024).

Trzecia strategia to automatyzacja etapów przeglądu, takich jak selekcja artykułów, ekstrakcja danych czy aktualizacja ustaleń w trybie dynamicznego przeglądu. W literaturze podkreśla się, że takie wsparcie zwiększa wydajność i może sprzyjać wyższej jakości metodycznej, o ile nad całym procesem czuwa badacz, odpowiedzialność za treść i wnioski pozostaje po stronie człowieka (Tomczyk, Brüggemann i Vrontis, 2024).

Czwarta strategia dotyczy mapowania wiedzy. Oprócz klasycznego science mapping można sięgnąć po metody, które skupiają się na zmiennych i ich relacjach, co lepiej wspiera formułowanie problemów badawczych i identyfikację luk badawczych w literaturze. Taki kierunek rozwija podejście variable science mapping i wskazuje, kiedy warto łączyć mapowanie z innymi technikami przeglądu (Tomczyk, Brüggemann i Paul, 2024).

Piąta strategia to świadome, transparentne i etyczne użycie narzędzi. Wytyczne instytucjonalne i środowiskowe są zgodne, należy jasno ujawniać, jaką rolę odegrała AI, zapewniać weryfikowalność i dbać o dobrostan badanych oraz integralność procesu. W praktyce oznacza to dokumentowanie użycia narzędzi, opis ról i ograniczeń modeli, a także

przygotowanie informacji dla recenzentów i czytelników (European Commission, 2025; Cheng, Calhoun i Reedy, 2025).

Szosta strategia akcentuje wątek twórczej kooperacji. Publikacje i projekty wygenerowanych książek pokazują, jak AI może stać się interaktywnym partnerem dialogu, inspirując koncepcje i wspierając iteracyjne pisanie, przy zachowaniu pełnej widoczności roli narzędzi i kontroli autora (Przegalińska i Triantoro, 2024).

Wreszcie, współpraca z AI powinna być osadzona w szerszej strategii zespołu i instytucji. Traktowanie AI jako stałego składnika strategii a nie dodatku ułatwia budowę kompetencji, wybór narzędzi i ustalenie standardów jakości oraz ujawniania użycia. Ten sposób myślenia wzmacnia przewagę konkurencyjną i odporność procesów badawczych (Przegalińska i Jemielniak, 2023; Beshr, Ateeq, Ateeq i Alaghbari, 2025).

7.3. Przykładowe ścieżki pracy z AI

Ścieżka 1 - Ustalenie zasad i odpowiedzialności - zaczynamy od wyboru narzędzi i spisania zasad ich użycia: jakie zadania wspiera AI, jak będzie dokumentowana jej rola oraz jak zapewniamy rzetelność i możliwość audytu. Wytyczne zalecają prostą notę o użyciu AI w manuskrypcie i w materiałach uzupełniających, wraz z opisem etapów, na których narzędzia wspierały pracę (European Commission, 2025; Cheng, Calhoun i Reedy, 2025).

Ścieżka 2 - Od pytania badawczego do planu - AI może pełnić rolę trenera koncepcyjnego: pomagać doprecyzować pytanie, wskazać możliwe operacjonalizacje i warianty projektu. Współpraca jest najcenniejsza, gdy pozostaje dialogiem - badacz formułuje kierunek, a narzędzie podsuwa hipotezy, szkice struktury i przykłady, które następnie są krytycznie oceniane i modyfikowane (Przegalińska i Triantoro, 2024; Haber, Przegalińska i Jemielniak, 2025).

Ścieżka 3 - Przegląd literatury łączony - najpierw wykonujemy zapytania w Scopusie i Web of Science, aby zbudować rdzeń przeglądu. Następnie prosimy narzędzia AI o wskazanie brakujących wątków i unikatowych pozycji. Z badań porównawczych wynika, że takie połączenie podnosi kompletność przeglądu: bazy tradycyjne zapewniają trafność i jakość, a systemy AI zwiększają szanse na odkrycie materiałów spoza typowych ścieżek (Tomczyk, Brüggemann, Mergner i Petrescu, 2024a; Tomczyk, Brüggemann, Mergner i Petrescu, 2024b).

Ścieżka 4 - Mapowanie pojęć i zmiennych - na etapie syntezy można wykorzystać mapowanie relacji między zmiennymi i teoriami, aby przejść od opisu słów kluczowych do obrazu zależności w polu badawczym. Takie mapowanie wspiera zarówno formułowanie

wkładu teoretycznego, jak i planowanie badań empirycznych (Tomczyk, Brüggemann i Paul, 2024).

Ścieżka 5 - Automatyzacja selekcji i ekstrakcji - do odsiewu rekordów, pozyskiwania metadanych i wyciągów z artykułów stosujemy narzędzia automatyzujące powtarzalne prace. Warto tworzyć dynamiczne przeglądy aktualizowane wraz z napływem nowych publikacji. Nad wynikami czuwają autorzy, którzy weryfikują trafność i korygują błędy modeli (Tomczyk, Brüggemann i Vrontis, 2024).

Ścieżka 6 - Pisanie, redakcja i kontrola jakości - generatywna AI pomaga wygładzać styl, skracać fragmenty i porządkować strukturę argumentacji, ale ostateczne brzmienie, fakty i interpretacje należą do autora. Dobrym standardem jest dodatkowa weryfikacja cytowań i skan antyplagiatowy, zwłaszcza kiedy narzędzia streszczają teksty zewnętrzne (Haber, Przeglasińska i Jemielniak, 2025).

Ścieżka 7 - Ujawnienie użycia i przygotowanie do recenzji - w manuskrypcie zamieszczamy krótką informację o roli AI, dołączamy szczegóły w aneksie i upewniamy się, że praktyki zgodne są z polityką czasopisma oraz instytucji. Rekomendacje instytucjonalne podkreślają znaczenie jasności i możliwości odtworzenia procesu, a także przypominają, by AI nie przypisywać autorstwa (European Commission, 2025; Cheng, Calhoun i Reedy, 2025; Haber, Przeglasińska i Jemielniak, 2025).

Ścieżka 8 - Osadzenie w strategii zespołu i wydziału - uzgadniamy standardy pracy z AI: listę narzędzi, minimalne wymagania dokumentacyjne i plan rozwoju kompetencji. Takie podejście, znane z zarządzania strategicznego, zwiększa spójność i odporność procesów badawczych, a także ułatwia zgodność z wartościami akademickimi (Przeglasińska i Jemielniak, 2023; Beshr, Ateeq, Ateeq i Alaghbari, 2025).

7.4. Checklisty weryfikacji (fakty, referencje, plagiat)

Rzetelna weryfikacja treści na końcu pracy badawczej to najszybszy sposób na uniknięcie recenzji odrzuconej za podstawowe błędy, takie jak: niespójne twierdzenia, błędne cytowania i niezamierzone zapożyczenia cudzych sformułowań. Dobrze ułożona lista kontrolna usprawnia ten etap, bo prowadzi krok po kroku przez trzy najbardziej wrażliwe obszary: fakty, referencje i oryginalność tekstu, a dopiero na końcu dodaje warstwę formalną. Poniżej znajduje się wprowadzenie do każdej z tych sfer oraz konkretne listy do sprawdzenia, które warto zastosować przed wysyłką manuskryptu.

1. Fakty i spójność wyводу

Na tym etapie chodzi o zgodność twierdzeń z danymi, logiką i materiałem dowodowym, który rzeczywiście był analizowany.

Checklista - fakty i spójność:

- Czy każde kluczowe twierdzenie ma w tekście jasne oparcie: w wynikach, tabelach, wykresach lub zacytowanych źródłach?
- Czy liczby są konsekwentne w całym manuskrypcie (w streszczeniu, wynikach, tabelach, opisie próby, wnioskach)?
- Czy podana jest nie tylko istotność, ale także wielkość efektu i przedziały ufności tam, gdzie to zasadne?
- Czy definicje pojęć i zmiennych są stałe, nie zmieniają znaczenia między sekcjami?
- Czy wnioski nie wykraczają poza to, co naprawdę wynika z metody i danych?
- Czy rysunki i tabele są numerowane, opisane i odsyłane w tekście; czy podpisy pozwalają zrozumieć je bez sięgania do głównego tekstu?
- Jeśli były wykorzystywane z narzędzi AI do streszczeń lub szkicu fragmentów: czy każdą informację zweryfikowano w pierwotnych źródłach?

2. Referencje i cytowania

Dobra bibliografia to nie tylko kompletność, ale i precyzja: źródło musi istnieć, pasować do tezy i być prawidłowo opisane.

Checklista - referencje:

- Czy cytowane jest źródło pierwotne zawsze, gdy to możliwe (zamiast wtórnych omówień)?
- Czy wszystkie pozycje z listy bibliograficznej występują w tekście i odwrotnie, żadnych braków i duplikatów?
- Czy każdy rekord ma poprawne dane: autorów, rok, tytuł, nazwy czasopism, numer tomu/zeszytu, strony, DOI lub URL?
- Czy styl i interpunkcja są spójne z wytycznymi (na przykład APA) w całym manuskrypcie?
- Czy nie ma nadmiernego samocytowania ani wąskiego doboru literatury, uwzględnione są prace reprezentujące różne stanowiska?
- Czy cytowania preprintów są aktualizowane na wersje recenzowane, jeśli już się ukazały?
- Jeśli AI podpowiadała źródła: czy każde było sprawdzone ręcznie (istnienie, treść, zgodność z tezą)?
- Czy przypisy w tekście odpowiadają liście bibliografii?

3. Plagiat, autoplagiat i oryginalność sformułowań

Oryginalność dotyczy nie tylko danych, ale też sposobu ujęcia. Weryfikować należy więc zarówno zapożyczenia literalne, jak i nadmierne parafrazy.

Checklista - oryginalność i etyka tekstu:

- Czy każdy cytat dosłowny jest ujęty w cudzysłów, z numerem strony i pełnym odniesieniem?
- Czy parafrazy rzeczywiście przetwarzają myśl cudzą własnymi słowami i z odwołaniem do autora (zamiast kosmetycznych podmian)?
- Czy nie powielono znaczących fragmentów własnych, wcześniej opublikowanych tekstów bez oznaczenia i uzasadnienia (autoplagiat)?
- Czy nie rozdrobniono wyników na kilka bardzo podobnych publikacji bez wyraźnie odrębnych pytań badawczych?
- Czy manuskrypt przeszedł przez przynajmniej jedno narzędzie detekcji podobieństwa tekstu i przejrane zostały wyniki jakościowo (kontekst dopasowań)?
- Jeśli AI wspierała redakcję: czy efekt końcowy nie zawiera nieoznaczonych, szablonowych fraz powszechnie generowanych przez modele oraz czy zachowuje indywidualny styl autora tekstu?

4. Transparentność korzystania z AI

Jeśli w procesie użyto narzędzi AI, do szkicu, porządkowania bibliografii, parafraz czy streszczeń, to warto to wskazać.

Checklista - przejrzystość użycia AI:

- Czy w sekcji Metodyka lub Oświadczenia”zwięźle wskazano, do czego AI wykorzystano i jak weryfikowano wyniki?
- Czy zadbano o brak danych wrażliwych w promptach i o zgodność z polityką czasopisma i uczelni?
- Czy potwierdzono, że modele nie są traktowane jako współautorzy, cała odpowiedzialność merytoryczna spoczywa na autorach?

5. Ostatni przegląd przed złożeniem manuskryptu

Krótką, techniczną pętlą pozwala uniknąć drobnych wpadek, które irytują recenzentów.

Checklista - kontrola końcowa:

- Tytuł, abstrakt i słowa kluczowe są spójne z treścią i terminologią używaną w dyscyplinie.
- Pliki uzupełniające (dane, skrypty, prerejestracje) są dostępne i opisane; linki działają.

- Oświadczenia: wkład autorów (CRediT), finansowanie, konflikty interesów, zgody etyczne - kompletne i konsekwentne.
- Format manuskryptu, rysunków i tabel zgodny z wytycznymi czasopisma; metadane w menedżerze bibliografii zsynchronizowane z tekstem.
- Wszyscy współautorzy widzieli wersję do wysyłki i zaakceptowali treść.

Podsumowanie rozdziału 7

AI może znacząco ułatwić proces tworzenia publikacji, zwiększając efektywność pracy i wspierając twórczość badacza. Jednocześnie jej użycie wymaga krytycznej weryfikacji wyników, aby uniknąć ryzyk związanych z halucynacjami, błędnymi cytowaniami czy utratą spójności argumentacji. Kluczowe pozostaje transparentne ujawnianie roli AI, zachowanie pełnej odpowiedzialności po stronie autora oraz traktowanie modeli nie jako zastępstwa, lecz jako narzędzia wspierającego. Dzięki temu sztuczna inteligencja może stać się wartościowym elementem warsztatu publikacyjnego, wzmacniając proces badawczy bez naruszania zasad rzetelności naukowej.

8. Automatyzacja w pracy badawczej

Automatyzacja stanowi jeden z kluczowych kierunków wykorzystania sztucznej inteligencji w pracy naukowej. Dzięki narzędziom typu no-code oraz integracji modeli językowych z popularnymi środowiskami, takimi jak arkusze kalkulacyjne czy menedżery bibliografii, badacze mogą znacząco usprawnić powtarzalne i czasochłonne zadania. Rozdział ten prezentuje przykłady praktycznych rozwiązań, które pozwalają przejść od ręcznego wykonywania analiz do zautomatyzowanych przepływów pracy, zapewniających większą efektywność i lepszą kontrolę nad danymi.

8.1. Koncepcja „Sheets → AI → raport”

Niniejszy rozdział prezentuje sposób zbudowania powtarzalnego procesu przetwarzania tekstu w arkuszu kalkulacyjnym, bez pisania kodu. Fundament stanowi Google Sheets oraz dodatek integrujący model językowy na przykład GPT for Sheets™ and Docs, który udostępnia funkcje wywoływane bezpośrednio w komórkach arkusza. Dodatek tego typu umożliwia generowanie odpowiedzi modeli AI na podstawie danych w arkuszu i jest dostępny w Google Workspace Marketplace.

Arkusz pełni rolę panelu sterowania. Każdy wiersz reprezentuje publikację lub inny rekord źródłowy. W kolumnach gromadzone są dane surowe na przykład abstrakty, a obok

umieszczane są formuły dodatku, które proszą model o ściśle zdefiniowane rezultaty: etykiety klasyfikacyjne, wyodrębnione pola na przykład metoda, próba lub zwięzłe streszczenia.

Przygotowanie środowiska

Instalacja: dodatek integrujący AI należy zainstalować z Marketplace i nadać mu uprawnienia do pracy w wybranych arkuszach.

Klucz API: do działania wymagany jest klucz API pozyskany z platformy dostawcy modelu, na przykład z OpenAI. Klucze API tworzy się i zarządza nimi w panelu OpenAI (sekcje API keys na koncie lub w projekcie). Klucz służy do uwierzytelniania wywołań i powinien pozostać poufny. Zalecane jest przechowywanie go wyłącznie w ustawieniach dodatku lub w bezpiecznym magazynie, bez udostępniania w treści arkusza.

Dobre praktyki dotyczące API

- Klucz należy traktować jako tajny; nie umieszczać go w komórkach arkusza ani w materiałach współdzielonych.
- Wskazane jest tworzenie osobnych kluczy dla projektów i ich okresowa rotacja.
- W miarę możliwości należy ograniczać uprawnienia i monitorować zużycie.

Projekt arkusza

Aby rozpocząć pracę z dodatkiem GPT for Sheets™ and Docs, należy zaprojektować arkusz w sposób umożliwiający systematyczne przechowywanie danych oraz wyników generowanych przez AI. Struktura arkusza powinna być prosta, ale jednocześnie przygotowana na różne typy informacji badawczych.

Zalecana jest prosta struktura:

- ID - identyfikator rekordu,
- Tytuł, Rok, Źródło,
- Abstrakt - główne pole tekstowe,
- kolumny wynikowe: Klasyfikacja, Metoda, Próba, Kontekst, Wynik główny, Streszczenie,
- kolumny meta: Data uruchomienia, Model, Wersja promptu.

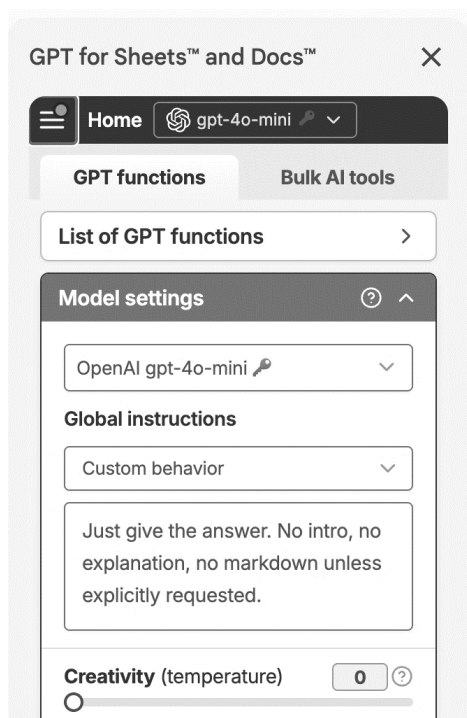
Zadanie: Należy utworzyć arkusz z wymienionymi kolumnami i wprowadzić 3 przykładowe rekordy (abstrakty). Następnie należy dodać puste kolumny na wyniki AI i metadane.

8.2. Klasyfikacja binarna abstraktów

Kolejnym krokiem jest przeprowadzenie klasyfikacji binarnej abstraktów na podstawie zadanego kryterium. Dzięki odpowiednio przygotowanemu promptowi możliwe jest szybkie uzyskanie powtarzalnych wyników w formie prostych odpowiedzi.

```
=GPT("Czy podany abstrakt dotyczy problematyki wykorzystania technologii blockchain w budowaniu lojalności klientów? Oznacz odpowiednio: tak, nie lub 'nie jestem w stanie określić'. Jeśli nie jesteś pewien, zwróć dokładnie 'nie jestem w stanie określić'."; L2)
```

Wskazówka. Dla stabilności wyników zalecane jest ustawienie możliwie niskiej kreatywności (temperatury jak przedstawiono na rysunku 8.1), aby uzyskiwać powtarzalne rezultaty.



Rys. 8.1. Panel GPT for Sheets™ and Docs z ustawieniami parametru temperatura

Źródło: opracowanie własne

Prompt - szablon klasyfikacji

Rola: recenzent. *Zadanie:* ocena spełnienia kryterium [tu definicja].

Format: wyłącznie tak, nie, nie jestem w stanie określić.

Materiał: L2.

Zadanie: Należy przeprowadzić klasyfikację dla wszystkich 3 abstraktów i ocenić ręcznie, czy zostały prawidłowo zaklasyfikowane.

8.3. Ekstrakcja informacji

W badaniach często istotne jest wyodrębnienie kluczowych elementów z abstraktów, takich jak metoda, próba, kontekst czy wynik główny. Automatyzacja tego procesu w arkuszu pozwala znacząco przyspieszyć analizę dużych zbiorów publikacji.

```
=GPT("Z poniższego abstraktu wyodrębnij i zwróć w JEDNYM wierszu cztery informacje w kolejności: metody; próba; kontekst; główny wynik. Jeśli brak informacji, wpisz 'brak'. Abstrakt: " ; L2)
```

Na rysunku 8.2 przedstawiono przykładowy zestaw abstraktów z wprowadzoną formułą pod działanie dodatku GPT for Sheets™ and Docs, ilustruje sposób, w jaki model może automatycznie uzupełniać kolumny arkusza.

▼ | *fx* =GPT("Z poniższego abstraktu wyodrębnij i zwróć w JEDNYM wierszu cztery informacje w kolejności: metody; próba; kontekst;

A	B	C	D	E	F
1	The impact of AI on customer experience Artificial intel	brak; brak; marketing; poprawa doświadczenia klienta			
2	Applications of Artificial Intelligence (AI) in Marketing M	brak; brak; marketing; przegląd podstawowy AI w zarządzaniu marketingiem			
3	How Artificial Intelligence Is Impacting Marketing? Artif	brak; brak; marketing; AI wpływa na działania, umiejętności i wydajność menedżerów marketingu.			
5	The impact of artificial intelligence on improving efficien	analiza literatury; brak; usługi; poprawa efektywności i zadowolenia klientów			

Rys. 8.2. Przykładowy zestaw abstraktów z wprowadzoną formułą pod działanie dodatku GPT for Sheets™ and Docs

Źródło: opracowanie własne

Zadanie: Należy przeprowadzić ekstrakcję dla wszystkich przykładowych rekordów i ocenić ręcznie, czy informacje nie zostały dopisane wbrew treści abstraktu. W razie potrzeby należy zaostrzyć sformułowania w *prompcie* (na przykład „nie uzupełniaj braków”).

8.4. Streszczenie robocze

Oprócz klasyfikacji i ekstrakcji informacji, możliwe jest również automatyczne generowanie krótkich streszczeń. Funkcja ta przydaje się zwłaszcza w analizach porównawczych lub przygotowaniu przeglądów literatury.

```
=GPT("Streść w 2 zdaniach (maks. 60 słów), styl naukowy, bez rekomendacji i bez źródeł. Zachowaj język oryginału. Abstrakt: " ; L2)
```

Kiedy nie streszczać? W analizach wrażliwych na styl oryginału (na przykład badania dyskursu) streszczenia powinny służyć wyłącznie orientacji, a nie jako materiał do kodowania.

8.5. Walidacja i replikowalność

Zalecane jest przygotowanie krótkiego codebooka kategorii. Następnie warto przeprowadzić kalibrację na losowej próbce (na przykład 20-30 rekordów ocenionych ręcznie) i dopiero potem zastosować formuły masowo. Dla replikowalności dobrze jest zapisywać obok

wyników: nazwę modelu, wersję *promptu*, datę/znacznik czasu i parametry (na przykład temperatura). Historia wersji arkusza ułatwia audyt zmian.

Prompt - kalibracja (few-shot)

„Oto trzy przykłady klasyfikacji (abstrakt → etykieta). Na ich podstawie oceń nowy przypadek. Odpowiadaj wyłącznie tak, nie, nie jestem w stanie określić. *Nowy abstrakt:*”; L2

8.6. Prywatność i zgodność

Do formuł nie powinny trafiać dane wrażliwe ani informacje objęte zobowiązaniami etycznymi lub prawnymi. W przypadku materiału empirycznego (ankiety, wywiady) wskazana jest anonimizacja. W dokumentach końcowych zaleca się przejrzystość: oznaczenie fragmentów, w których wykorzystano narzędzia AI, oraz opis sposobu walidacji.

8.7. Najczęstsze problemy i rozwiązania

Podczas pracy z dodatkiem mogą pojawić się różne trudności techniczne. Warto znać typowe źródła błędów oraz sposoby ich rozwiązywania, aby praca z arkuszem przebiegała sprawniej.

- #ERROR! lub przerwane działanie - możliwe przekroczenie limitów dodatku; włączenie przełącznika RUN i partiami uruchamianie przeliczenia zwykle stabilizują pracę.
- Separator argumentów - w polskiej lokalizacji obowiązują średniki ;.
- Zmienność wyników - zaleca się niższą kreatywność (temperatura ~0) i doprecyzowanie formatu.

8.8. Ćwiczenie: mini-przegląd literatury

Na zakończenie niniejszego rozdziału proponowane jest praktyczne ćwiczenie. Dzięki niemu można przećwiczyć wszystkie omówione wcześniej funkcjonalności, a także samodzielnie ocenić potencjalne ograniczenia i źródła błędów.

Należy przygotować 30 abstraktów z wybranego tematu i wprowadzić je do arkusza. Kolejno: (1) przeprowadzić klasyfikację binarną zgodnie z kryterium własnym; (2) wyodrębnić cztery informacje (3) wygenerować krótsze streszczenia. Na końcu wskazane jest dopisanie dwóch zdań o ograniczeniach procedury i potencjalnych źródłach błędów.

Podsumowanie rozdziału 3

Omówione w rozdziale przykłady pokazują, że automatyzacja z wykorzystaniem AI nie wymaga zaawansowanych umiejętności programistycznych, a może realnie odciążyć badacza w zakresie organizacji i analizy danych. Automatyzacja nie zastępuje krytycznego myślenia ani interpretacji, ale umożliwia skupienie się na aspektach twórczych i koncepcyjnych, które decydują o wartości naukowej badań.

Dodatek AI do Google Sheets przekształca arkusz w laboratorium tekstowe do klasyfikacji, ekstrakcji i streszczania. Kluczem są: precyzyjne *prompty* z wymuszonym formatem, proste bezpieczniki (IF, RUN), systematyczna walidacja oraz dokumentowanie parametrów pracy. Klucz API należy pozyskać z zaufanego dostawcy (na przykład OpenAI) i przechowywać zgodnie z dobrymi praktykami bezpieczeństwa.

9. Detekcja treści generowanych przez AI

Rozwój narzędzi generatywnej sztucznej inteligencji sprawił, że rozróżnienie między treściami tworzonymi przez człowieka, a tekstami generowanymi przez modele językowe stało się istotnym wyzwaniem dla nauki i edukacji. W odpowiedzi powstały systemy detekcji AI, które analizują właściwości języka, styl pisanie czy statystyczne wzorce tokenów w celu określenia prawdopodobieństwa, że dany tekst został wygenerowany automatycznie. Rozdział ten omawia najważniejsze narzędzia detekcyjne, ich ograniczenia oraz konsekwencje stosowania w praktyce akademickiej.

9.1. Detektory AI

Większość narzędzi do wykrywania tekstów generowanych przez modele językowe (LLM) wykorzystuje statystyczne cechy języka: perplexity (przewidywalność kolejnych tokenów), burstiness (zmiennosc dystrybucji tokenów), stylometrię oraz log-prawdopodobieństwo sekwencji wyliczane przez model odniesienia. Klasyczne podejścia pokazują, że teksty maszynowe częściej zalegają w obszarach o niższej entropii i większej przewidywalności, co można wizualizować lub klasyfikować (Gehrmann, Strobel, Rush, 2019; Mitchell, Lee, Khazatsky, Manning, Finn, 2023).

Alternatywną klasą metod jest watermarking, czyli celowe, niewidoczne dla człowieka znakowanie generatu podczas tworzenia, które później ułatwia identyfikację jego pochodzenia. W praktyce komercyjnej narzędzia łączą te sygnały w modele klasyfikacyjne i raportują wynik w postaci prawdopodobieństwa lub udziału AI w tekście (Kirchenbauer, Geiping, Wen, Katz, Miers, Goldstein, 2023).

Narzędzie ZeroGPT udostępnia klasyfikację AI/Human oraz wskaźnik procentowy, bazując według deklaracji na miarach pokroju perplexity/burstiness oraz modelach uczenia maszynowego (zerogpt.com). W literaturze akademickiej narzędzie pojawia się w

przekrojowych porównaniach, gdzie bywa konkurencyjne dla starszych modeli na przykład GPT-3.5, ale jego skuteczność spada przy nowszych modelach lub tekstach modyfikowanych (Walters, 2023). Na rysunku 9.1 widać, że dla próbki tekstu wygenerowanego przez model GPT-5 Pro narzędzie oceniło AI GPT na 58,39%, co wskazuje na wynik niejednoznaczny.

Your Text is Most Likely AI/GPT generated



Rys. 9.1. Wynik ZeroGPT dla tekstu wygenerowanego przez model GPT-5 Pro

Źródło: badanie własne w narzędziu zerogpt.com

Natomiast na rysunku 9.2 analogiczna próbka tekstu wygenerowanego przez model GPT-4o została oznaczona jako 100% AI/GPT.

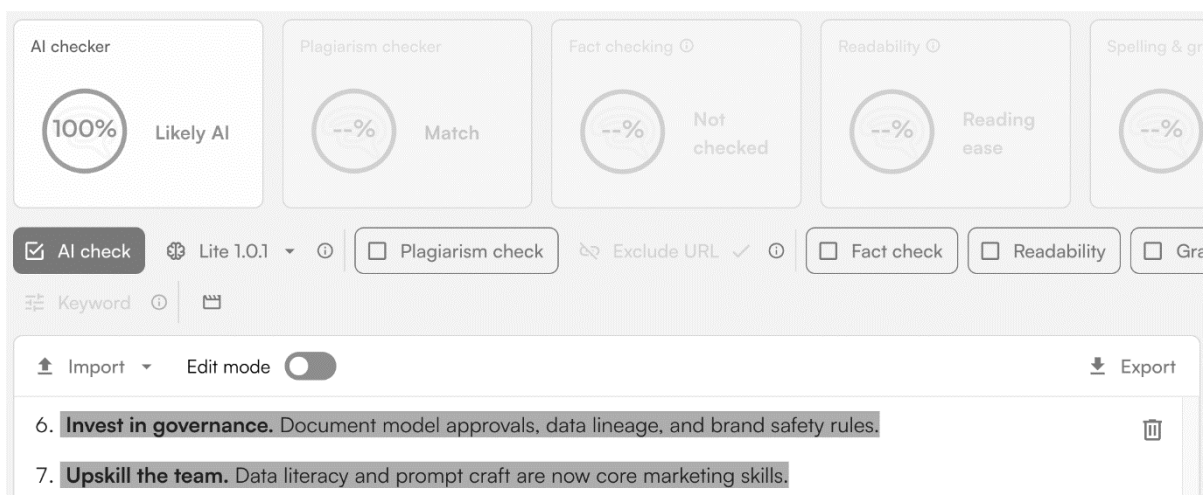
Your Text is AI/GPT Generated



Rys. 9.2. Wynik ZeroGPT dla tekstu wygenerowanego przez model GPT-4o

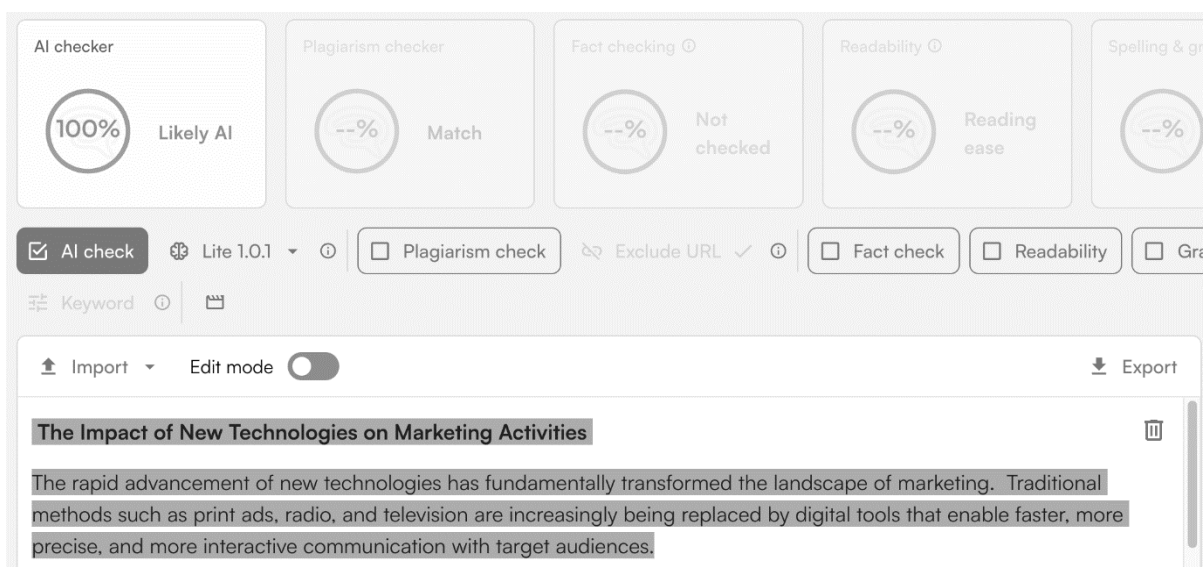
Źródło: badanie własne w narzędziu zerogpt.com

Originality.ai łączy detekcję AI z plagiatem i raportowaniem na poziomie zdań (Originality.ai). W niezależnym badaniu porównawczym 16 detektorów (126 esejów: GPT-3.5, GPT-4 oraz prace studenckie) Originality.ai, obok Copyleaks i Turnitin, osiągnął wysoką trafność dla wszystkich trzech zbiorów, podczas gdy wiele innych narzędzi miało problemy zwłaszcza z tekstami z GPT-4 (Walters, 2023). Na rysunku 9.3 i rysunku 9.4 pokazano wyniki 100% Likely AI dla obu próbek wygenerowanych tekstów przez modele GPT-5 Pro i GPT-4o, co ilustruje zdecydowaną klasyfikację tego detektora w testach poglądowych.



Rys. 9.3. Wynik Originality.ai dla tekstu wygenerowanego przez model GPT-5 Pro

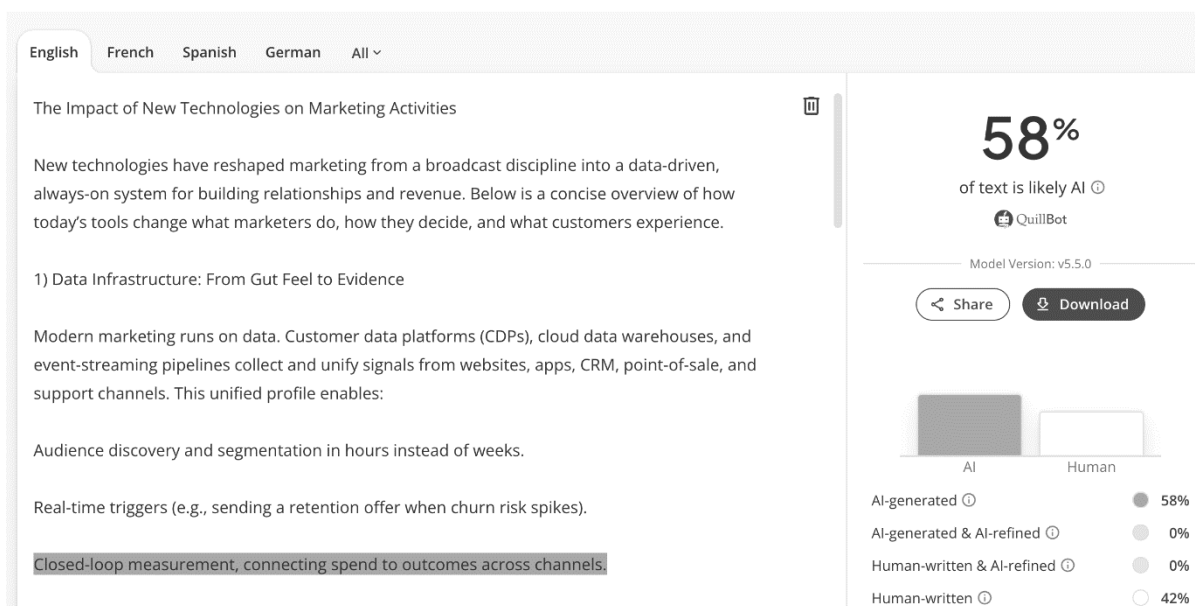
Źródło: badanie własne w narzędziu Originality.ai



Rys. 9.4. Wynik Originality.ai dla tekstu wygenerowanego przez model GPT-4o

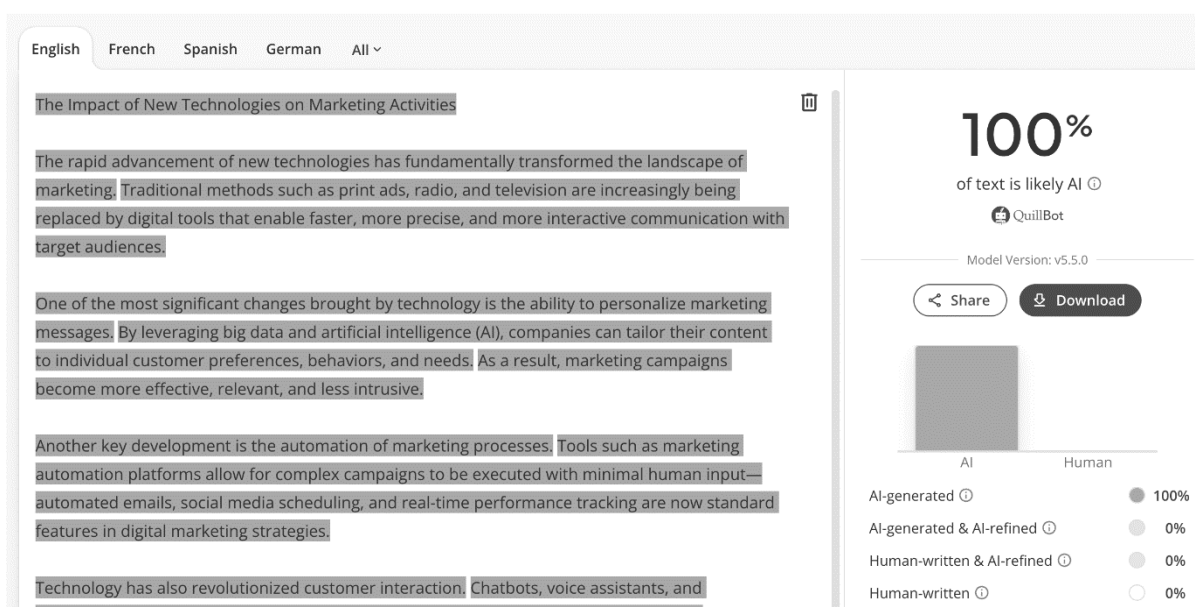
Źródło: badanie własne w narzędziu Originality.ai

QuillBot pierwotnie znany z parafrazowania oferuje również wykrywacz AI, który raportuje udział tekstu likely AI (quillbot.com). Publicznie dostępne opisy działania mówią o analizie cech językowych i porównywaniu wzorców z danymi treningowymi, w literaturze naukowej narzędzie jest oceniane rzadziej i raporty o jego skuteczności mają charakter mieszany. Na rysunku 9.5 detektor wskazał 58% likely AI zatem wynik częściowy, natomiast na rysunku 9.6 zidentyfikował 100% likely AI.



Rys. 9.5. Wynik QuillBot AI Detector dla tekstu wygenerowanego przez model GPT-5 Pro

Źródło: badanie własne w narzędziu quillbot.com



Rys. 9.6. Wynik QuillBot AI Detector dla tekstu wygenerowanego przez model GPT-4o

Źródło: badanie własne w narzędziu quillbot.com

Dwa często cytowane podejścia badawcze to GLTR, czyli wizualizacja pozycji słów w rozkładach prawdopodobieństwa dla wskaźników top-k/nucleus oraz DetectGPT - wykorzystanie krzywizny prawdopodobieństwa warunkowego „curvature” do zero-shotowej detekcji. Z kolei watermarking (Kirchenbauer i in., 2023) nie analizuje stylu tekstu, lecz wykrywa sygnał wbudowany przy generowaniu, wymaga to jednak współpracy twórcy modelu i bywa podatne na parafrazowanie (Gehrmann i in., 2019; Mitchell i in., 2023; Kirchenbauer i in., 2023).

9.2. Falszywe alarmy i granice detekcji

Detektory źle generalizują poza rozkłady treningowe (inne domeny, gatunki, rejestry). Producent OpenAI wyraźnie informował, że jego (obecnie już wycofany) klasyfikator ma niski poziom trafności, jest bardzo zawodny dla krótkich tekstów poniżej 1000 znaków, gorzej działa w językach innych niż angielski i nie powinien służyć jako narzędzie decyzyjne pierwszego wyboru. To ostrzeżenie pozostaje aktualne dla większości bieżących detektorów (OpenAI, 2023).

Szereg badań pokazuje, że parafrazowanie nawet lekkie, automatyczne lub ręczne dramatycznie obniża skuteczność wykrywania (w tym DetectGPT, GPTZero i klasyfikatorów opartych na watermarkach). Metody rekurencyjnego parafrazowania (recursive paraphrasing) potrafią obniżyć pole pod krzywą ROC (Area Under the Receiver Operating Characteristic Curve) nawet o rzędy wielkości, przy znikomym spadku jakości językowej (Krishna, Song, Karpinska, Wieting, Iyyer, 2023; Sadasivan, Kumar, Balasubramanian, Wang, Feizi, 2023).

W badaniu Liang i in. (2023) detektory systematycznie myliły prace osób piszących po angielsku jako języku obcym, przypisując im autorstwo AI częściej niż native speakerom. Oznacza to realne ryzyko karania stylu odmiennego od danych treningowych (Liang, Yuksekgonul, Mao, Wu, Zou, 2023).

W ewaluacji 16 narzędzi część detektorów odróżniała GPT-3.5 od człowieka, lecz była generalnie nieskuteczna wobec GPT-4, najlepsi Copyleaks, Turnitin, Originality.ai utrzymali jednak wysoką trafność na obu zbiorach. Skuteczność jest zatem wrażliwa na wersje modelu i stale się zmienia (Walters, 2023).

Wynik X% AI jest heurystyczny, zależny od długości próbki i przyjętego progu decyzyjnego, ten sam tekst może zostać oceniony od niejednoznaczny do 100% AI przez różne narzędzia (rysunki 1–6 dla tej samej pary próbek tekstów wygenerowanych przez modele GPT-5 Pro i GPT-4o). Użytkownik powinien traktować te wskaźniki jako wskazówki, a nie dowód (OpenAI, 2023).

W świetle badań recenzowanych najsilniejsze potwierdzenie w niniejszej trójce ma Originality.ai (stabilnie wysokie wyniki w porównaniu 16 narzędzi, również na próbkach z GPT-4). ZeroGPT bywa skuteczny, ale w niezależnych ewaluacjach jego trafność jest bardziej zmienna, szczególnie przy nowszych modelach i tekstach redagowanych. QuillBot AI Detector ma najsłabsze pokrycie w literaturze naukowej, dostępne są głównie opisy marketingowe i testy dziennikarskie, dlatego jego wyniki należy interpretować szczególnie ostrożnie (Walters, 2023).

Jeśli detektor w ogóle ma być używany w procesie dydaktycznym lub redakcyjnym, to jako pomocniczy i najlepiej w zestawie co najmniej dwóch narzędzi, z dodatkową weryfikacją procesu powstawania tekstu (historia wersji, notatki, szkice, logika cytowań).

Rekomendacje dla młodych naukowców i doktorantów:

- Nie opierać decyzji wyłącznie na detektorze. Nawet producenci narzędzi i badacze podkreślają ograniczoną niezawodność i wysokie ryzyko fałszywych alarmów w krótkich lub przetworzonych tekstach. Należy zabezpieczać się dowodami procesu takimi jak: wersje robocze, notatki, repozytoria kodu (OpenAI, 2023).
- Weryfikować wyniki wieloma źródłami. Łączyć co najmniej dwa detektory z oceną merytoryczną: spójność metodologii, poprawność cytowań, kontrola plagiatu.
- Uważać na stronniczość (bias AI). W przypadku autorów piszących w języku obcym lub prac z nietypową nomenklaturą należy uwzględniać możliwość stronniczości modelu (Liang, 2023)
- Rozumieć techniczne ograniczenia. Detekcja statystyczna $\neq$ rozstrzygnięcie autorstwa, watermarking wymaga wsparcia po stronie generatora i również bywa obchodzony przez parafrazowanie (Kirchenbauer, 2023).

Podsumowanie rozdziału 9

Detektory AI są użyteczne jako wskaźniki ryzyka, ale nie zapewniają dowodu autorstwa. W zaprezentowanych testach próbek tekstów (rysunki 1-6) widać rozbieżności między narzędziami, szczególnie przy tekście wygenerowanym przez model GPT-5 Pro, a literatura naukowa wykazuje, że parafrazowanie, redakcja i różnice populacyjne autorów znacząco zwiększają liczbę fałszywych wyników.

Z trzech omawianych narzędzi Originality.ai ma najsilniejsze potwierdzenie empiryczne w badaniach porównawczych, mimo to każdy wynik musi być interpretowany ostrożnie i łączony z oceną procesu powstawania tekstu (Sadasivan, 2023; Walters, 2023).

10. Tworzenie i aktualizacja materiałów dydaktycznych

Generatywna sztuczna inteligencja może być też narzędziem, które będzie wspierać przygotowywanie materiałów dydaktycznych. Jej cechami, które to umożliwiają, a które do tej pory rozpoznaliśmy jest szybkość, analiza i synteza olbrzymich ilości danych, możliwość dostosowania do potrzeb użytkownika. Tym co jest nowym elementem, który może być szczególnie istotny w pracy dydaktycznej, jest multimodalność (Smith, i in., 2025).

Dotyczy to możliwości korzystania i generowania zasobów tekstowych, dźwiękowych oraz graficznych. W przypadku pracy dydaktycznej pozwala to na wykorzystanie różnych form reprezentacji treści celem dotarcia do studentów kanałem, który jest dla w danym momencie najbardziej dostępny. „Badania nad cyfrowym komponowaniem multimodalnym w przeważającej mierze podkreślają potencjał zwiększania zaangażowania uczniów, potwierdzania tożsamości i wspierania innowacyjnego tworzenia znaczeń”. (Smith, i in., 2025). Multimodalność genAI zwiększa sprawczość dydaktyczną, ponieważ pozwala tworzyć i analizować treści w formatach bliższych praktyce akademickiej (tekst–obraz–audio–video), a także wspiera refleksję nad nauczaniem na bazie danych wideo. (Imran, i in. 2024). Dlatego korzystając z obecnych narzędzi genAI bardzo łatwo nam poruszać się i transformować treści w kanały tekst-audio, audio-tekst, wizualizacje, symulacje, a nawet platformy i narzędzie wspierające AR i VR (Hemminki-Reijonen, i in., 2025; Lampropoulos, 2025).

Podobnie jak przy każdym aspekcie pracy młodego naukowca, mającej wnosić wartość, również w przypadku działań dydaktycznych, niezbędnym elementem jest opracowanie celu jakiemu służyć będzie wykorzystanie danej pomocy lub zasobu. Dlatego pomimo łatwości uzyskania prezentacji, kart pracy, opisów przypadków czy innych bardziej złożonych technologicznie materiałów istotą staje się włączenie tych aspektów w proces projektowania lub poszerzania środowiska uczenia się studentów. odejście to przypomina framework ARCHED (Responsible, Collaborative, Human-centered Education Instructional Design) (Li i in., 2025), który eksponuje konieczność precyzyjnego określenia celów dydaktycznych, zgodnych z taksonomią Blooma, już na etapie projektowania materiału wspieranego AI, z zachowaniem centralnej roli nauczyciela jako głównego decydenta procesu dydaktycznego.

10.1. Tworzenie sylabusów i studiów przypadku

Przygotowanie treści sylabusów czy studiów przypadku do pracy dydaktycznej ze studentami często związane jest ze żmudnym działaniem, które wymaga dużego nakładu czasu. Jednak obecnie mając do dyspozycji bazę wcześniej przygotowanych sylabusów czy studiów przypadków, wraz z wytycznymi dotyczącymi ich zawartości, proces ich przygotowania można sobie ułatwić dzięki wykorzystaniu generatywnej sztucznej inteligencji.

Kluczową rolę odgrywa tutaj praca na bazie zebranych i zweryfikowanych danych: wcześniejszych sylabusów, studiów przypadku oraz wytycznych instytucjonalnych, które stanowią fundament jakości i spójności. Na ich podstawie AI może proponować modyfikacje istniejących dokumentów, przygotowywać wstępne szkice lub tworzyć nowe treści dla poszczególnych kategorii (np. efekty uczenia się, metody oceny, literatura). Całość, podobnie

jak w przypadku pracy w innym zakresie, przebiega w oparciu o powtarzalny cykl, a w kolejnych etapach można doskonalić poszczególne części sylabusów oraz podejmować decyzję na jakim etapie zakończyć współpracę z chatbotem.

Przykładowy prompt (możliwy do modyfikacji do własnych potrzeb) może wyglądać tak:

Otrzymujesz zbiór dokumentów zawierających wcześniej przygotowane sylabusy oraz wytyczne instytucjonalne dotyczące kursów akademickich. Na ich podstawie przygotuj nowy sylabus dla przedmiotu [wstaw nazwę].

Struktura sylabusa powinna obejmować:

- nazwę kursu i opis ogólny,*
- cele i efekty uczenia się zgodne z [wskaz źródło wymagań],*
- plan tygodniowy,*
- metody dydaktyczne,*
- metody oceny,*
- literaturę podstawową i uzupełniającą.*

Przygotuj dokument w taki sposób, aby możliwa była jego modyfikacja: oznacz sekcje, w których wykładowca może wprowadzić własne uzupełnienia (np.: [Do uzupełnienia przez prowadzącego]).

Na końcu zaproponuj alternatywną wersję sylabusa w wariantcie projektowym lub opartym na studiach przypadku.

10.2. Tworzenie quizów, zadań i prezentacji

GenAI może również wspierać nas w przygotowywaniu zadań, quizów i prezentacji. Bardzo często w znanych do tej pory narzędziach twórcy usług dodają funkcjonalności wspierane przez modele AI, które umożliwiają automatyzację i wsparcie w generowaniu treści. W przypadku zadań i quizów możemy do generowania zasobów i wkładu używać chatbotów, oraz wyspecjalizowanych narzędzi.

Tabela 10.1. Przykładowe narzędzia do tworzenia quizów

Narzędzie	Dostępność	Kluczowe funkcjonalności wspierane przez AI
www.questgen.ai		

Questgen	Bezpłatne	Generuje testy wielokrotnego wyboru, prawda/fałsz, kreowanie z różnych źródeł (tekst, PDF, web, obraz, wideo). Eksport w wielu formatach (Moodle, PDF).
www.revisely.com		
Revisely AI Quiz Maker	Bezpłatne	Przekształca notatki, książki, PDF-y, PowerPoint w quizy; wspiera cały proces (tworzenie, edycja, ocena).
questionwell.org		
QuestionWell	Bezpłatne	Tworzy, analizuje, personalizuje i eksportuje materiały zgodne z wymaganiami programowymi.
www.quillionz.com		
Quillionz	Bezpłatne	Generuje pytania i streszczenia z tekstu w czterech krokach; zaprojektowane jako narzędzie wspomagające nauczanie.

Źródło: Opracowanie własne

Dodatkowo warto pamiętać, że te i podobne narzędzia można często integrować z LMS (*Learning Management System*) takimi jako Google Workspace, czy Office 365. Również w nich znajdują się narzędzia pozwalające tworzyć zadania, monitorować przebieg ich realizacji oraz korzystać ze wsparcia chatbotów takich jak Copilot czy Gemini. Platforme te mogą być też pomocą przy kreowaniu prezentacji. Dodatkowo warto korzystać również z narzędzi typu Gamma (www.gamma.app) czy Canva (www.canva.com). Ta ostatnia jest potężnym narzędziem projektowania graficznego i wizualnego, wyposażonym w moduł AI wspierający projektowanie dokumentów, reprezentacji wizualnych, materiałów filmowych treści interaktywnych.

10.3. Metaprompty i refleksja nad uczeniem się

W procesie dydaktycznym można również wykorzystać chatboty, które będą wspierały pracę studentów. W tym obszarze według prof. Joanny Mytnik ważne jest budowanie przestrzeni do refleksji nad korzyściami i trudnościami związanymi z sięganiem po te narzędzia. Rekomenduje ona sięganie po karty pracy, z pomocą których uczący porządkują swoje doświadczenia, rozwijając swoje kompetencje krytycznego myślenia i uczenia się (Mytnik, 2025).

Dodatkową cenną propozycją są metaprompty, których celem jest takie poinstruowanie chatbota, co do jego stylu komunikacji i zadań, aby mógł być asystentem procesu uczenia się i pełnił dedykowane zadania.

Przykładowy metaprompt, gotowy do dostosowywania do własnych potrzeb znajduje się poniżej (CNE PG, 2024):

„Jesteś wspierającym nauczycielem-tutorem, który pomaga osobie uczącej się zrozumieć koncepcje, wyjaśnia zagadnienia i zadaje pytania. Zaczynij od przedstawienia się osobie uczącej się jako jego AI-Tutor, który chętnie odpowie na wszelkie pytania i zapytaj, jak się masz zwracać do osoby uczącej się. Zadawaj tylko jedno pytanie na raz. Najpierw zapytaj, czego chciałby się dowiedzieć. Poczekaj na odpowiedź. Następnie zapytaj o poziom znajomości tematu, stan zaawansowania wiedzy. Poczekaj na odpowiedź. Następnie zapytaj, co już wie na wybrany temat. Poczekaj na odpowiedź. Biorąc pod uwagę te informacje, pomóż osobie uczącej się zrozumieć temat podając wyjaśnienia, przykłady, analogie. Powinny one być dostosowane do poziomu i jego wcześniejszej wiedzy lub tego, co już wie na dany temat. Podaj osobie uczącej się wyjaśnienia, przykłady i analogie dotyczące koncepcji, aby pomóc mu zrozumieć. Powinieneś prowadzić osobę uczącą się w sposób otwarty. Nie udzielaj natychmiastowych odpowiedzi lub rozwiązań problemów, ale pomóż mu wygenerować własne odpowiedzi poprzez zadawanie pytań naprowadzających. Poproś osobę uczącą się o wyjaśnienie jego sposobu myślenia. Jeśli ma trudności lub otrzymuje błędną odpowiedź, spróbuj poprosić go o wykonanie części zadania lub przypomnij mu o jego celu i daj mu wskazówkę. Jeśli osoba ucząca się poprawi się, pochwal i okaż radość. Jeśli osoba ucząca się ma trudności, zachęć go i daj mu kilka pomysłów do przemyślenia. Naciskając na osobę uczącą się w celu uzyskania informacji, spróbuj zakończyć swoje odpowiedzi pytaniem, aby student musiał nadal generować pomysły. Gdy osoba ucząca się wykaże odpowiedni poziom zrozumienia, biorąc pod uwagę jego poziom nauki, poproś ją o wyjaśnienie koncepcji własnymi słowami; to najlepszy sposób, aby pokazać, że coś wie lub poproś go o przykłady. Gdy osoba uczącej się wykaże, że zna pojęcie, możesz zakończyć rozmowę i powiedzieć mu, że jesteś tutaj, aby pomóc, jeśli ma dalsze pytania.”

Podsumowanie rozdziału 10

Generatywna sztuczna inteligencja otwiera nowe możliwości w przygotowywaniu i doskonaleniu materiałów dydaktycznych, umożliwiając szybkie tworzenie sylabusów, studiów przypadku, quizów, prezentacji czy treści multimedialnych. Jej zastosowanie pozwala zwiększyć różnorodność form nauczania, lepiej dopasować materiały do potrzeb studentów oraz zautomatyzować czasochłonne procesy. Jednocześnie niezbędne jest zachowanie krytycznej roli nauczyciela akademickiego – to on nadaje ostateczny kształt materiałom, zapewnia ich zgodność merytoryczną i spójność z celami dydaktycznymi. AI może być zatem traktowana jako wsparcie, które podnosi efektywność i atrakcyjność zajęć, ale nigdy nie zastępuje odpowiedzialności i kompetencji dydaktyka.

11. Najlepsze praktyki i rekomendacje dla młodych naukowców i doktorantów

Punktem wyjścia dotyczącym standardów praktyk młodych naukowców są dokumenty opracowane przez międzynarodowe instytucje. W niniejszym rozdziale sięgniemy po dwa z nich. Jeden to opublikowana w roku 2024 przez Radę Europy publikacja „Europejski kodeks postępowania na rzecz uczciwości w badaniach naukowych” (ang. „The European Code of Conduct for Research Integrity”). Drugi to przewodnik „Wytyczne dotyczące generatywnej sztucznej inteligencji w edukacji i badaniach naukowych” (ang. „Guidance for generative AI in education and research”) opracowany i opublikowany przez UNESCO w 2023 roku. Porządkują one i tworzą ramy dla pracy dydaktycznej i naukowej w sytuacji dynamicznego rozwoju sztucznej inteligencji.

11.1. Równowaga między efektywnością a rzetelnością

Raport "The European Code of Conduct for Research Integrity" (ALLEA, 2024) argumentuje, że rzetelność badań jest kluczowa dla utrzymania wiarygodności systemu badawczego i jego wyników. Obszar ten obejmuje wszystkie dziedziny życia naukowego i akademickiego, w tym różnych środowisk badawczych. W związku z tym dotyczy także osób, które na podstawie różnych uregulowań prawnych zaliczają się do grona młodych naukowców. Opracowany kodeks służy jako rama samoregulacji dla społeczności badawczej.

Najważniejsze zasady rzetelności badawczej to (ALLEA, 2024):

- niezawodność obejmująca jakość badań, na każdym etapie obejmującym planowanie, metodologię, analizę i wykorzystanie zasobów;
- uczciwość podczas opracowywania, prowadzenia, recenzowania, raportowania i komunikowania badań zgodnie z zasadami przejrzystości, sprawiedliwości i bezstronności;
- szacunek dla osób ze środowiska badawczego, uczestników badań, podmiotów badawczych, społeczeństwa, ekosystemów, dziedzictwa kulturowego i naturalnego;
- odpowiedzialność za badania na każdym etapie ich realizacji, od pomysłu do publikacji, również w zakresie zarządzania i organizacji, za szkolenia, nadzór i mentoring, a także za szerszy wpływ społeczny.

Kodeks poświęca szczególną uwagę roli szkolenia, nadzoru i mentoringu, co jest kluczowe dla rozwoju młodych naukowców i doktorantów. Wskazuje, że badacze na wczesnym etapie kariery i zatrudnieni na krótkoterminowych umowach mogą być szczególnie narażeni na

zagrożenia związane z presją efektywności, przez co rzetelność badań, w których są zaangażowani może być zagrożona. Kodeks mocno akcentuje aspekty, które zapewniają jakość, niezawodność i solidność wyników badań, co w dłuższej perspektywie przekłada się na rzeczywistą wartość, często kosztem szybkiej, ale powierzchownej "efektywności".

Kluczowe aspekty promujące rzetelność, które wspierają właściwą równowagę, to (ALLEA, 2024):

- nacisk kładziony jest na niezawodność w zapewnianiu jakości badań, odzwierciedloną w projekcie, metodologii, analizie i wykorzystaniu zasobów. Badania powinny być projektowane, przeprowadzane, analizowane i dokumentowane w staranny, przejrzysty i dobrze przemyślany sposób.
- naukowcy powinni raportować swoje wyniki i metody, w tym użycie zewnętrznych usług lub narzędzi AI, w sposób zgodny z przyjętymi normami dyscypliny i ułatwiający weryfikację lub powtórne użycie procedury tam, gdzie to możliwe.
- dostęp do danych powinien być tak otwarty, jak to możliwe, i tak zamknięty, jak to konieczne, zgodny z zasadami FAIR (dane możliwe do odczytania (*ang. Findable*), dostępne (*ang. Accessible*), interoperacyjne (*ang. Interoperable*) i do ponownego użytku (*ang. Reusable*);
- autorzy powinni być dokładni i uczciwi w komunikacji, a także przejrzysti w kwestii założeń, wartości, solidności dowodów, niepewności i luk w wiedzy. Ważne jest publikowanie również wyników negatywnych, które są tak samo istotne jak pozytywne.
- kodeks wyraźnie wymienia praktyki naruszające rzetelność, takie jak fabrykowanie, fałszowanie, plagiat, manipulowanie statystykami, ukrywanie użycia AI, "publikacje salami" (rozdrabnianie wyników) czy selektywne cytowanie. Unikanie tych praktyk jest fundamentalne dla zachowania rzetelności nad szybkim, ale nieszczerze uzyskanym wynikiem.

Treść raportu UNESCO jest spójna z powyższymi założeniami. Jego treść rozszerza temat rzetelności w badania naukowych o zagrożenia związane z pogłębianiem się cyfrowych nierówności, naruszenie praw autorskich, luk poznawczych (*ang. bias*), redukcji różnorodności i spłaszczenia opinii, nieprawidłowych danych i informacji, a także kosztów środowiskowych (UNESCO, 2023).

Wskazówki z wyżej przedstawionych dokumentów stanowią dobre podłoże, aby wskazać, że obszar nauki nie powinien opierać się tylko na aspektach związanych z wydajnością, którą narzędzia sztucznej inteligencji mogą wspierać. Większość wartość ma projektowanie przestrzeni do uwolnienia się od tempa i presji, które mogą również w przypadku stosowania cyfrowych rozwiązań obniżać jakość pracy naukowej i dydaktycznej (Berg, Seeber, 2016). Tym, co istotnie wpływa na jakość pracy naukowej, jest świadomość, że odpowiednio

długi okres czasu sprzyja dobrej organizacji projektów zespołowych i wspomaga proces refleksji oraz wartościowego zarządzania i realizowania zadań (NAS 2015).

11.2. Matryca ryzyka/korzyści dla różnych etapów badań

Koncepcja oceny i minimalizowania ryzyka jest silnie obecna w raporcie EU, szczególnie w sekcji poświęconej środkom ochronnym (*ang. safeguards*). Korzyści są związane z dążeniem do wiedzy i zwiększeniem zrozumienia świata oraz z zachowaniem wiarygodności badań.

Kluczowe punkty dotyczące ryzyka i korzyści:

- naukowcy powinni rozpoznawać i rozważać potencjalne szkody i ryzyka związane z badaniami i zastosowaniami oraz łagodzić możliwe negatywne skutki;
- badacze, instytucje i organizacje muszą przestrzegać odpowiednich kodeksów, wytycznych i przepisów. Powinni traktować uczestników i podmioty badań (ludzi, zwierzęta, dziedzictwo kulturowe, środowisko) z szacunkiem i ostrożnością, zgodnie z przepisami prawnymi i zasadami etycznymi;
- należy należycie dbać o zdrowie, bezpieczeństwo i dobrostan społeczności, współpracowników i innych osób związanych z badaniami;
- niezastosowanie się do dobrych praktyk badawczych szkodzi procesom badawczym, niszczy relacje między naukowcami, podważa zaufanie do badań i ich wiarygodność, marnuje zasoby i może narażać uczestników badań, podmioty, użytkowników, społeczeństwo lub środowisko na niepotrzebne szkody.

Bazując na istniejących w literaturze wskazówkach opracowaliśmy matrycę ryzyka i korzyści wspierającą proces kontroli i podejmowania decyzji w procesie realizacji badań. Może ona dla młodego naukowca pełnić narzędzie orientacyjne: pokazuje, jak na różnych etapach badań generatywna AI może wspierać lub osłabiać rzetelność procesu. Jednocześnie może też być narzędziem refleksji, poprzez porównywania własnych przykładów z obowiązującymi standardami. Sekcja decyzji pozwala posłużyć się wytycznymi wskazującymi kierunek konkretnych działań. Poniższa tabela jest propozycją i nie zastępuje krytycznej oceny badacza, lecz ją wspiera. W tabeli zastosowano podział na etapy badań (koncepcja, projekt/protokół, dane, analiza/kod, pisanie, upowszechnianie, zarządzanie ryzykiem) stanowi syntezę klasycznych ujęć procesu badawczego (Creswell & Creswell, 2018).

Tabela 11.1. Analiza ryzyka i korzyści dla wykorzystania genAI

Etap badań	Potencjalne korzyści (AI)	Główne ryzyka	Kontrola	Próg decyzyjny
------------	---------------------------	---------------	----------	----------------

1. Koncepcja / mapowanie literatury	szybki przegląd, streszczenia, mapy tematów	halucynacje, błędne cytowania, luki poznawcze (ang. Bias)	weryfikacja ręczna źródeł, korzystanie ze zweryfikowanych baz; notatnik audytowy dokumentujący decyzje, działania i użyte zasoby	Tylko gdy >80% źródeł ma DOI; inaczej warunkowo. (UNESCO, 2023)
2. Projekt / protokół	triangulacja metod, checklisty PRISMA/CONSO RT	„fałszywy autorytet” modeli, niedopasowanie metod	zgodność z listami kontrolnymi PRISMA/CONSORT;	Tylko po zgodności z checklistą. (Page i in., 2021; CONSORT, 2025)
3. Dane	szablony planów oparte na FAIR, anonimizacja wstępna	prywatność, licencje, nieuprawnione dane udostępnione dla modelu	repozytoria instytucjonalne; plan zgodny z FAIR	Tylko gdy spełnione FAIR i warunki przetwarzania danych. (Wilkinson i in., 2016)
4. Analiza / kod	generowanie szkiców kodu, testów, komentarzy	błędy obliczeń, dług replikacyjny (ang. <i>reproducibility debt</i>)	archiwizacja różnych wersji, testy jednostkowe, notatnik działań i obliczeń, archiwizacja kodu/danych	Tylko z testami i archiwizowaniem wersji. (Nosek i in., 2015)
5. Pisanie / redakcja	propozycje struktury, przyspieszenie szkiców, redakcja językowa	plagiat, „zmyślone” źródła, styl nieakademicki	ujawnienie użycia, pełna weryfikacja cytowań, zakaz traktowania AI jako autora	Tylko z ujawnieniem i pełną weryfikacją (COPE, 2023)
6. Upowszechni anie / grafika	streszczenia popularnonaukow e, wizualizacja treści	prawa autorskie, różnice kontekstowe	sprawdzenie polityki wydawcy używaj własnych/wolnych zasobów	Tylko zgodnie z polityką wydawcy. (Nature, 2023; Springer Nature, 2023)
7. Zarządzanie ryzykiem	profilowanie sieci i zależności ryzyka i kontroli	niedoszacowanie obszarów i jego skutków w różnych obszarach	nadzorowanie, mapowanie, mierzenie, zarządzenie kryzysem	Zgodnie z macierzą ryzyka (NIST, 2023)

Źródło: Opracowanie własne

Oczywiście w przypadku młodych naukowców nieocenione jest wsparcie mentorów, promotorów czy środowiska badawczego w którym proces rozwoju się odbywa. Wsparcie

doświadczonych badaczy jest ciągle istotnym i kluczowym elementem wspierania młodych naukowców w procesie rozwoju warsztatu dydaktycznego i naukowego.

11.3. Plan rozwoju kompetencji cyfrowych

Kompetencje cyfrowe to szeroki zakres łączący funkcje wiedzy, umiejętności i postaw, które wykraczają poza wąskie rozumienie w postaci myślenia komputacyjnego czy technicznej sprawności w zakresie wykorzystywania cyfrowych urządzeń. Uznawane są za jedną z ośmiu kluczowych kompetencji życiowych dla uczenia się przez całe życie. (European Commission, 2006). To zestaw umiejętności pozwalających na efektywne wykorzystywanie technologii w celu optymalizacji codziennego życia. (Zhao, i in. 2021) Europejska Komisja opisuje je jako pewne, krytyczne i odpowiedzialne korzystanie z technologii społeczeństwa informacyjnego w pracy, rozrywce i edukacji (Education and Training, 2019).

W obszarze wiedzy młody naukowiec powinien posiadać rozeznanie jak technologie cyfrowe (w tym różne rozwiązania wspierane sztuczną inteligencją) mogą wspierać realizację zadań dydaktycznych i naukowych, komunikacyjność, kreatywność, a także być świadomym ich ograniczeń, skutków ryzyka. Wiąże się to także, ze znajomością aspektów etycznych i regulacyjnych, wielokrotnie wspomnianych w tej publikacji. (Education and Training, 2019) W obszarze umiejętności osoby będące na początku drogi w pracy naukowej powinny potrafić wykorzystywać cyfrowe technologie do współpracy, kreatywności, prowadzenia projektów i realizacji zadań dydaktycznych. Umiejętności te obejmują zdolność do korzystania, dostępu, filtrowania, tworzenia, programowania i udostępniania treści cyfrowych (Education and Training, 2019).

W obszarze postaw młodzi naukowcy powinni kultywować i rozwijać refleksyjną i krytyczną postawę, oraz otwartość i zorientowanie na wartościowe możliwości korzystania z cyfrowych rozwiązań. Nieodłączna jest potrzeba wzmacniania odpowiedzialnego i bezpiecznego korzystania z tego typu rozwiązań. (Education and Training, 2019).

Dla nauczycieli i kadry akademickiej, kompetencje cyfrowe to nie tylko krytyczne i odpowiedzialne korzystanie z zasobów ICT, ale również inne istotne umiejętności dotyczące oceny, współpracy i udzielania informacji zwrotnej studentom w ramach zwyczajowej praktyki edukacyjnej w świecie cyfrowym (dos Santos, i in., 2023).

W badaniach naukowych, umiejętność radzenia sobie z danymi w sposób zaplanowany i bezpieczny, a także ich świadome i etycznie odpowiednie wykorzystanie, staje się kluczowe w niemal wszystkich dziedzinach. Ważne jest odpowiednie archiwizowanie, publikowanie i

możliwość ponownego wykorzystania danych, co jest istotne dla integralności i wiarygodności badań oraz ich odtwarzalności. (dos Santos, i in., 2023).

Europejskie Ramy Kompetencji Cyfrowych zostały opracowane przez Wspólne Centrum Badawcze (ang. *Joint Research Center*) Komisji Europejskiej. Przedstawiają one charakterystykę opisującą sześć poziomów rozwoju od poziomu podstawowego do pionierskiego. Nazwy poszczególnych poziomów zostały powiązane z nomenklaturą wywodzącą się z Europejskiego Systemu Opisu Kształcenia Językowego (CERF), stosując oznaczenia od A1 do C2. Proces przechodzenia przez kolejne etapy rozwoju to ciągły proces obejmujący również korzystanie z genAI. Jego strukturę przedstawia tabela 18 (Vuorikari, i in., 2022).

Tabela 11.2. Model rozwoju kompetencji cyfrowych

Poziom (CEFR)	Opis kryteriów (DigCompEdu)	Działania służące osiągnięciu poziomu	Poziom korzystania z genAI
Nowicjusz (A1) (ang. <i>Newcomer</i>)	Podstawowa świadomość możliwości technologii, nieregularne i fragmentaryczne wykorzystanie, głównie do administracji i prostych zadań dydaktycznych.	Uczestnictwo w szkoleniach wprowadzających; korzystanie z prostych narzędzi (np. prezentacje, e- platformy).	Eksperymentalne użycie AI do prostych podpowiedzi (np. poprawa języka w e- mailu, wygenerowanie listy tematów).
Odkrywca (A2) (ang. <i>Explorer</i>)	Zainteresowanie eksploracją narzędzi cyfrowych, stosowanie ich w wybranych obszarach, działanie niesystematyczne.	Włączenie narzędzi cyfrowych do niektórych procesów dydaktycznych lub projektów badawczych; korzystania z wiedzy i umiejętności praktyków.	Użycie genAI do wyszukiwania literatury, tworzenia krótkich podsumowań, draftów quizów lub sylabusów.
Integrator (B1) (ang. <i>Integrator</i>)	Eksperymentowanie z technologiami w różnych kontekstach i przy osiąganiu różnych celów, integracja z praktyką dydaktyczną i naukową.	Tworzenie własnych zasobów dydaktycznych online, stosowanie blended learning, planowanie zajęć w oparciu o systemy zarządzania nauczaniem (LMS).	Systematyczne użycie genAI jako narzędzia wspomagającego tworzenie treści dydaktycznych i analiz danych.
Ekspert (B2) (ang. <i>Ekspert</i>)	Krytyczne i świadome korzystanie z szerokiej gamy technologii, dopasowywanie narzędzi do konkretnych celów i sytuacji.	Udział w projektach wdrażających technologie cyfrowe na uczelni, mentoring dla innych pracowników.	Krytyczne wykorzystanie genAI do redakcji tekstów, analizy jakościowej, przygotowania wizualizacji i symulacji.

Lider, (C1) (ang. <i>Leader</i>)	Kompleksowe i strategiczne podejście, inspiracja dla innych, ciągłe doskonalenie i poszukiwanie nowych i wydajnych rozwiązań.	Projektowanie szkoleń z zakresu kompetencji cyfrowych, publikowanie dobrych praktyk, kierowanie zespołami dydaktycznymi i badawczymi.	Rozwijanie metapromptów i strategii pracy z genAI, prowadzenie szkoleń dla doktorantów i współpracowników. Integrowanie genAI do zespołów i projektów badawczych.
Pionier, (C2) (ang. <i>Pioneer</i>)	Innowacyjne, eksperymentalne podejście, kreowanie nowych modeli dydaktycznych i technologicznych, wdrażanie i upowszechnianie unikalnych rozwiązań w środowisku akademickim.	Inicjowanie projektów badawczych nt. innowacji cyfrowych, współpraca międzynarodowa, opracowywanie rekomendacji.	Eksperymentowanie z multimodalną AI, tworzenie własnych narzędzi AI, integracja AI w badaniach i dydaktyce jako element nowej metodologii.

Źródło: opracowanie własne na podstawie DigCompEdu

Podsumowanie rozdziału 11

Niniejszy szósty rozdział podkreśla, że kluczowym wyzwaniem dla młodych naukowców jest zachowanie równowagi między efektywnością, jaką daje sztuczna inteligencja, a rzetelnością badań. Międzynarodowe wytyczne (UNESCO i Rady Europy) wskazują, że presja produktywności nie może prowadzić do obniżenia jakości pracy, fundamentem pozostają uczciwość, przejrzystość, dokładność i odpowiedzialność za wyniki. Narzędzia AI mogą wspierać każdy etap procesu badawczego, jednak wymagają krytycznej oceny źródeł, weryfikacji danych, ochrony prywatności oraz zgodności z normami etycznymi i prawnymi. Pomocą w tym zakresie może być matryca ryzyka i korzyści, która ułatwia podejmowanie odpowiedzialnych decyzji.

Równolegle niniejszy rozdział akcentuje znaczenie rozwijania kompetencji cyfrowych nie tylko technicznych, ale także krytycznego myślenia, komunikacji, kreatywności i współpracy. Umiejętności te powinny ewoluować stopniowo: od podstawowej orientacji w narzędziach, przez systematyczne wykorzystywanie AI w dydaktyce i badaniach, aż po innowacyjne modele pracy naukowej. Całość wskazuje, że odpowiedzialna integracja AI w warsztacie badawczym wymaga nie tylko sprawności technicznej, ale także etycznej refleksji, transparentności i wsparcia mentorów oraz instytucji.

Zakończenie

Podczas pisania niniejszej monografii autorzy mieli w głowie myśl, aby przygotować użyteczny przewodnik, po który będą sięgać doktoranci i młodzi naukowcy. Zebrano w nim definicje, perspektywy, modele oraz praktyczne przykłady, mogące służyć całej społeczności na wczesnym etapie kariery badawczej. Zakłada się, że korzystanie z zawartości książki ułatwi uporządkowane i holistyczne wykorzystywanie możliwości generatywnej sztucznej inteligencji, a równocześnie pomoże rozwijać świadomość jej ograniczeń. W zgodzie z planami działań badawczych oraz uniwersalnymi ramami etycznymi zaleca się ustawiczny rozwój i poszerzanie kompetencji, także poza zakresem przedstawionym w tej publikacji. Tempo zmian i niepewność co do przyszłych kierunków rozwoju technologii nie pozwalają jednoznacznie przesądzić, jak narzędzia AI będą oddziaływać na instytucje naukowe. Niezmiennie aktualne pozostają jednak wartości wskazywane w książce, stanowiące pomost między teraźniejszością a nadchodzącą przyszłością.

W części pierwszej uporządkowano pojęcia i zarysowano kontekst teoretyczny, opisując miejsce AI, zwłaszcza modeli generatywnych w metodologii współczesnych badań. Pokazano synergię człowieka i systemów obliczeniowych w perspektywie uczenia się, odwołując się do modeli cyklicznej pracy badawczej, oraz zidentyfikowano ograniczenia: halucynacje, stronniczości, problemy z wyjaśnialnością i ryzyka prawno-etyczne. Wskazano równocześnie na zmiany metodyczne (dynamiczne przeglądy, integracje RAG, rosnące wymagania transparentności) oraz na zróżnicowane, lecz zbieżne co do zasady stanowiska wydawnictw naukowych dotyczące jawności i odpowiedzialności autorów. Całość osadzono w aktualnym stanie praktyk i regulacji, podkreślając konieczność stałej weryfikacji polityk i standardów.

Druga część miała charakter przewodnika operacyjnego. Zaproponowano tu ramy inżynierii promptów i szablony pracy etapowej (od ekstrakcji informacji po weryfikację wyników), przedstawiono scenariusze wykorzystania AI w pisaniu i redakcji, a także przykłady automatyzacji „Sheets → AI → raport” oraz dobre praktyki w tworzeniu materiałów dydaktycznych. Zebrane checklisty weryfikacji treści i referencji akcentują, że narzędzia generatywne mogą przyspieszać pracę, ale nie znoszą obowiązku krytycznej oceny, replikowalności procedur i rzetelności bibliograficznej. Wskazano zarazem granice detekcji treści AI, rekomendując traktowanie takich wyników jako sygnałów pomocniczych, nie rozstrzygnięć.

W całym opracowaniu przyjęto stanowisko pragmatyczne: AI stanowi element infrastruktury badawczej, który wzmacnia analizę, syntezę i komunikację naukową, lecz nie

zastępuje ludzkiej odpowiedzialności za konceptualizację, wybory metodologiczne i interpretację. Dlatego nacisk położono na praktyki „human-in-the-loop”, czyli dokumentowanie użycia narzędzi, włączanie weryfikacji krzyżowej i jawne ujawnianie roli modeli w tekście. Takie podejście łączy dążenie do efektywności z wymogami jakości i integralności badań.

Rekomendacje sformułowane na gruncie analiz i przykładów sprowadzają się do trzech zasad ogólnych. Po pierwsze, priorytet ma przejrzystość: należy opisywać dane, decyzje i narzędzia w sposób umożliwiający audyt i odtworzenie ścieżki dowodowej. Po drugie, kluczowa jest adekwatność metodyczna: AI ma wspierać zadania, w których przynosi mierzalną wartość, a nie generować pozór precyzji. Po trzecie, potrzebna jest długofalowa polityka kompetencyjna, rozwój umiejętności wykorzystania AI w pisaniu, kompetencji transwersalnych i praktyk zespołowych powinien być planowany i ewaluowany tak, jak inne elementy warsztatu badawczego.

Niniejsza monografia nie wyczerpuje podjętego tematu. Odnotowano dynamiczność zarówno samych modeli, jak i stanowisk wydawniczych oraz uwarunkowań prawnych, zaleca się zatem regularną aktualizację procedur zgodnie z wymogami dyscypliny, czasopism i instytucji oraz rozwijanie praktyk otwartej nauki, licencjonowania i ochrony danych. Przyjęte tu rozwiązania (szczególnie te dotyczące przeglądów, automatyzacji i pracy dydaktycznej) mają charakter adaptowalnych wzorców, które można dopasowywać do specyfiki obszaru badawczego.

Wobec niepewności technologicznej stabilnym punktem odniesienia pozostają wartości i standardy rzetelności: wiarygodność, uczciwość, szacunek i odpowiedzialność. To one nakazują utrzymywać człowieka w roli decydenta, a AI traktować jako zasób wymagający nadzoru, kontroli jakości i jawności. Taki paradygmat współpracy umożliwia wykorzystywanie potencjału narzędzi generatywnych bez uszczerbku dla istoty pracy naukowej, poszukiwania prawdy, które jest zarazem procesem poznawczym, społecznym i etycznym.

Monografia może zatem pełnić funkcję kompasu: porządkuje podstawy, proponuje metody i oferuje wzorce działania, nie zastępując krytycznego namysłu i autonomii badawczej. Włączenie jej zaleceń do codziennej praktyki sprzyja budowaniu środowiska badań i kształcenia, w którym generatywna AI staje się narzędziem zwiększającym sprawczość i zasięg nauki, a nie substytutem ludzkiego osądu.

Spis tabel i rysunków

Tabele

2.1. Integracja modeli Resnicka, Kolba, badań naukowych i roli genAI	29
2.2. Poziomy kompetencji w modelu SAMR.....	33
5.1. Zestawienie kluczowych zasad polityki Elsevier wobec generatywnej AI	47
5.2. Zestawienie kluczowych zasad polityki Springer Nature wobec generatywnej AI.....	48
5.3. Zestawienie kluczowych zasad polityki Taylor & Francis wobec generatywnej AI ...	49
5.4. Zestawienie kluczowych zasad polityki Wiley wobec generatywnej AI	50
5.5. Zestawienie kluczowych zasad polityki SAGE wobec generatywnej AI.....	50
5.6. Zestawienie kluczowych zasad polityki IEEE wobec generatywnej AI	51
5.7. Zestawienie kluczowych zasad polityki Emerald wobec generatywnej AI.....	52
5.8. Zestawienie do praw autorskich treści wygenerowanych z wykorzystaniem genAI	54
6.1. Składowe procesu promptowania	62
6.2. Składowe promptowania wg White i inni Nazwa i klasyfikacja	62
6.3. Przykładowa struktura prompta.....	63
6.4. Wsparcie twórców genAI w zakresie tworzenia promptów	67
10.1. Przykładowe narzędzia do tworzenia quizów	102
11.1. Analiza ryzyka i korzyści dla wykorzystania genAI.....	108
11.2. Model rozwoju kompetencji cyfrowych.....	110

Rysunki

Rys. 1.1. Liczba publikacji wymieniających narzędzia AI w bazie Google Scholar (2019 - 2025)..	22
Rys. 1.2. Liczba publikacji wymieniających narzędzia AI w bazach Web of Science i Scopus (2019-2025)	23
Rys. 2.1. Spirala kreatywnego uczenia się	26
Rys. 2.2. Cykl uczenia się opartego na doświadczeniu	26
Rys. 2.3. Diagram dwunastu kompetencji w zakresie generatywnej AI.....	30
Rys. 7.1. Ustawienia Asystenta scite.ai panel „Assistant Settings” z najważniejszymi przełącznikami.....	76
Rys. 7.2. Przykładowy wynik odpowiedzi Asystenta scite.ai na pytanie badawcze	77
Rys. 7.3. Lista źródeł i karta Search Strategy	78
Rys. 7.4. Upload papers - dodawanie PDF/BIB/RIS do prywatnej biblioteki Elicit	79
Rys. 7.5. Manage Columns - tworzenie lub wybór kolumn do ekstrakcji danych	80
Rys. 7.6. Edit custom column - pole Instructions do formułowania własnych promptów.....	81
Rys. 7.7 Menu kolumny - edycja, duplikowanie, filtrowanie, zapisywanie jako preset.....	81
Rys. 7.8. Tabela ekstrakcji - przykładowe wypełnienie kolumn Main findings i Methodology	82
Rys. 8.1. Panel GPT for Sheets TM and Docs z ustawieniami parametru temperatura	91
Rys. 8.2. Przykładowy zestaw abstraktów z wprowadzoną formułą pod działanie dodatku GPT for Sheets TM and Docs	92
Rys. 9.1. Wynik ZeroGPT dla tekstu wygenerowanego przez model GPT-5 Pro	95
Rys. 9.2. Wynik ZeroGPT dla tekstu wygenerowanego przez model GPT-4o	95
Rys. 9.3. Wynik Originality.ai dla tekstu wygenerowanego przez model GPT-5 Pro	96
Rys. 9.4. Wynik Originality.ai dla tekstu wygenerowanego przez model GPT-4o.....	96
Rys. 9.5. Wynik QuillBot AI Detector dla tekstu wygenerowanego przez model GPT-5 Pro	97
Rys. 9.6. Wynik QuillBot AI Detector dla tekstu wygenerowanego przez model GPT-4o	97

Literatura

1. Aljamaan, F., Temsah, M.-H., Altamimi, I., AlEyadhy, A., Jamal, A., Alhasan, K., Mesallam, T. A., Farahat, M., & Malki, K. H. (2024). *Reference Hallucination Score for Medical Artificial Intelligence Chatbots: Development and Usability Study*, JMIR Medical Informatics, 12, e54345, <https://doi.org/10.2196/54345> (dostęp: 20.07.2025).
2. Alkaissi, H., McFarlane, S. I. (2023), *Artificial hallucinations in ChatGPT: Implications in scientific writing*, Cureus, 15(2), e35179, <https://doi.org/10.7759/cureus.35179> (dostęp: 20.07.2025).
3. ALLEA All European Academies (2023) *The European Code of Conduct for Research Integrity*, Revised Edition, Berlin: ALLEA, <http://allea.org/wp-content/uploads/2023/06/European-Code-of-Conduct-Revised-Edition-2023.pdf> (dostęp: 11.07.2025).
4. Annapureddy, R., Fornaroli, A., Gatica-Pérez, D. (2025). *Generative AI Literacy: Twelve Defining Competencies*, *Digital Government: Research and Practice*, 1, s. 1–21
5. Anthropic (2023). *Claude terms of service*, <https://privacy.anthropic.com/en/collections/10672414-policies-terms-of-service> (dostęp: 07.07.2025).
6. Anthropic (2024). *Prompt engineering overview*. *Anthropic Documentation*, <https://docs.anthropic.com/en/docs/build-with-claude/prompt-engineering/overview> (dostęp: 07.07.2025).
7. Ateriya, N., Sonwani, N. S., Thakur, K. S., Kumar, A., Verma, S. K. (2025). *Exploring the ethical landscape of AI in academic writing*. *Egyptian Journal of Forensic Sciences*, 15, 36.
8. Athaluri, S. A., Manthena, S. V., Kesapragada, V. S. R. K. M., Yarlagaadda, V., Dave, T., Duddumpudi, R. T. S. (2023). *Exploring the boundaries of reality: Investigating the phenomenon of artificial intelligence hallucination in scientific writing through ChatGPT references*, Cureus, 15(4), e37432, <https://doi.org/10.7759/cureus.37432> (dostęp: 10.08.2025).
9. Banerjee, S., Agarwal A., Singla S. (2024). *LLMs will always hallucinate, and we need to live with this*. <https://arxiv.org/abs/2409.05746> (dostęp: 20.06.2025).
10. Barbierato, E., Gatti, A. (2024). *The Challenges of Machine Learning: A Critical Review*. *Electronics*, 13(2), 416, <https://doi.org/10.3390/electronics13020416> (dostęp: 20.06.2025).
11. Bennett, M. T. (2022). *Computable Artificial General Intelligence*, arXiv:2205.10513, <https://doi.org/10.48550/arXiv.2205.10513>, (dostęp: 13.07.2025).

12. Berg, M. and Seeber, B.K. (2016) *The slow professor: Challenging the culture of speed in the academy*, Toronto: University of Toronto Press.
13. Beshr, B., Ateeq, A., Ateeq, R. A., Alaghbari, M. A. (2025). *AI in Academia: Pros and Cons of Integrating Artificial Intelligence in Universities*, W: E. AlDhaen (red.), *Business Sustainability with Artificial Intelligence (AI): Challenges and Opportunities*, Springer.
14. Beshr, B., Ateeq, A., Ateeq, R. A., Alaghbari, M. A. (2025). *AI in academia: Pros and cons of using Artificial Intelligence. W: Business Sustainability with Artificial Intelligence (AI): Challenges and Opportunities* (Studies in Systems, Decision and Control, 566), Springer.
15. Bhattacharyya, M., Miller, V. M., Bhattacharyya, D., & Miller, L. E. (2023). *High rates of fabricated and inaccurate references in ChatGPT-generated medical content*, Cureus, 15(5), e39238. <https://doi.org/10.7759/cureus.39238>, (dostęp: 2.07.2025).
16. Bockting, C. L., van Dis, E. A. M., van Rooij, R., Zuidema, W., & Bollen, J. (2023). *Living guidelines for generative AI-Why scientists must oversee its use*, Nature, 622, 693–696.
17. Boiko, D. A., McKnight R., Kline B., Gomes G. (2023). *Autonomous chemical research with large language models*, Nature, 624(7992), 570–578, <https://doi.org/10.1038/s41586-023-06792-0>, (dostęp: 10.08.2025).
18. Bommasani, R., Hudson, D., Adeli, E., Altman, R., Arora, S., von Arx, S., Bernstein, M., Bohg, J., Bosselut, A., Brunskill, E., Brynjolfsson, E., Buch, S., Card, D., Castellon, R., Chatterji, N., Chen, A., Creel, K., Davis, J., Demszky, D., Donahue, C., Doumbouya, M., Durmus, E., Ermon, S., Etchemendy, J., Ethayarajh, K., Fei-Fei, L., Finn, C., Gale, T., Gillespie, L., Goel, K., Goodman, N., Grossman, S., Guha, N., Hashimoto, T., Henderson, P., Hewitt, J., Ho, D., Hong, J., Hsu, K., Huang, J., Icard, T., Jain, S., Jurafsky, D., Kalluri, P., Karamcheti, S., Keeling, G., Khani, F., Khattab, O., Koh, P., Krass, M., Krishna, R., Kuditipudi, R., Kumar, A., Ladhak, F., Lee, M., Lee, T., Leskovec, J., Levent, I., Li, X., Li, X., Ma, T., Malik, A., Manning, C., Mirchandani, S., Mitchell, E., Munyikwa, Z., Nair, S., Narayan, A., Narayanan, D., Newman, B., Nie, A., Niebles, J., Nilforoshan, H., Nyarko, J., Ogut, G., Orr, L., Papadimitriou, I., Park, J., Piech, C., Portelance, E., Potts, C., Raghunathan, A., Reich, R., Ren, H., Rong, F., Roohani, Y., Ruiz, C., Ryan, J., Ré, C., Sadigh, D., Sagawa, S., Santhanam, K., Shih, A., Srinivasan, K., Tamkin, A., Taori, R., Thomas, A., Tramèr, F., Wang, R., Wang, W., Wu, B., Wu, J., Wu, Y., Xie, S., Yasunaga, M., You, J., Zaharia, M., Zhang, M., Zhang, T., Zhang, X., Zhang, Y., Zheng, L., Zhou, K., Liang, P. (2021). *On the Opportunities and Risks of Foundation Models*, arXiv:2108.07258, <https://doi.org/10.48550/arXiv.2108.07258> (dostęp: 21.06.2025).

19. Borgdorff, H. (2012). *The Conflict of the Faculties: Perspectives on Artistic Research and Academia*,
<https://library.oapen.org/viewer/web/viewer.html?file=/bitstream/handle/20.500.12657/32887/595042.pdf?sequence=1&isAllowed=y> (dostęp: 11.07.2025).
20. Borgeaud, S., Mensch, A., Hoffmann, J., Cai, T., Rutherford, E., Millican, K., Van Den Driessche, G. B., Lespiau, J.-B., Damoc, B., Clark, A., De Las Casas, D., Guy, A., Menick, J., Ring, R., Hennigan, T., Huang, S., Maggiore, L., Jones, C., Cassirer, A., Brock, A., Paganini, M., Irving, G., Vinyals, O., Osindero, S., Simonyan, K., Rae, J., Elsen, E., Sifre, L. (2022). *Improving language models by retrieving from trillions of tokens. Proceedings of the 39th International Conference on Machine Learning (ICML)*, 2206–2240,
<https://arxiv.org/abs/2112.04426> (dostęp: 11.07.2025).
21. Bouhouita Guermech, S., Gogognon, P., BélislePipon, J.C. (2023). *Specific challenges posed by artificial intelligence in research ethics. Frontiers in Artificial Intelligence*, 6, 1149082,
<https://www.frontiersin.org/journals/artificial-intelligence/articles/10.3389/frai.2023.1149082/full> (dostęp: 11.07.2025).
22. Brnabic, A., Hess, L. M. (2021). Systematic literature review of machine learning methods used in clinical decision making. *BMC Medical Informatics and Decision Making*, 21, 62,
<https://doi.org/10.1186/s12911-021-01403-2> (dostęp: 13.07.2025).
23. Brown, T., Mann, B., Ryder, N., Subbiah, M., Kaplan, J., Dhariwal, P., Neelakantan, A., Shyam, P., Sastry, G., Askell, A., Agarwal, S., Herbert-Voss, A., Krueger, G., Henighan, T., Child, R., Ramesh, A., Ziegler, D., Wu, J., Winter, C., Hesse, C., Chen, M., Sigler, E., Litwin, M., Gray, S., Chess, B., Clark, J., Berner, C., McCandlish, S., Radford, A., Sutskever, I., Amodei, D. (2020). Language Models are FewShot Learners. *Advances in Neural Information Processing Systems*, 33, 1877–1901, (NeurIPS proceedings), <https://arxiv.org/abs/2005.14165> (dostęp: 18.08.2025).
24. Cancino, M., Panes, J. (2021). The impact of Google Translate on L2 writing quality measures: Evidence from Chilean EFL high school learners, *System*, 98, 102464,
<https://doi.org/10.1016/j.system.2021.102464> (dostęp: 13.07.2025).
25. Carlini, N., Tramèr, F., Wallace, E., Jagielski, M., Herbert-Voss, A., Lee, K., Roberts, A., Brown, T., Song, D., Erlingsson, Ú., Oprea, A., Raffel, C. (2021). *Extracting training data from large language models*, USENIX Security 2021,
<https://www.usenix.org/conference/usenixsecurity21/presentation/carlini-extracting> (dostęp: 12.06.2025)

26. CEDEFOP (2023) *Skills for Learning: Trend Focus*. Luxembourg: Publications Office of the European Union, <https://www.cedefop.europa.eu/en/tools/skills-intelligence/trend-focus/skills-learning>, (dostęp: 04.07.2025).
27. Centrum Nowoczesnej Edukacji PG (2023). *Wytyczne dotyczące wykorzystania narzędzi generatywnej AI*, Politechnika Gdańska, <https://cne.pg.edu.pl/genai/wytyczne>, (dostęp: 07.07.2025).
28. Chelli, M., Lemos, M., i in. (2024). *Hallucination rates and reference accuracy of ChatGPT and Bard for systematic reviews: Comparative analysis*, Journal of Medical Internet Research, 26, e53164, <https://doi.org/10.2196/53164>, (dostęp: 07.07.2025).
29. Cheng, A., Calhoun, A., Reedy, G. (2025). *Artificial intelligence-assisted academic writing: Recommendations for ethical use*, Advances in Simulation, 10, 22, <https://doi.org/10.1186/s41077025003506>, (dostęp: 14.07.2025).
30. Cheng, M., Calhoun, A. T., Reedy, S. (2025). *Artificial intelligence–assisted academic writing: Recommendations for ethical use*. Advances in Simulation, 10, 22, <https://advancesinsimulation.biomedcentral.com/articles/10.1186/s41077-025-00350-6> (dostęp: 22.06.2025).
31. Chowdhery, A., Narang, S., Devlin, J., Bosma, M., Mishra, G., Roberts, A., Barham, P., Chung, H. W., Sutton, C., Gehrmann, S., Schuh, P., Shi, K., Tsvyashchenko, S., Maynez, J., Rao, A., Barnes, P., Tay, Y., Shazeer, N., Prabhakaran, V., Reif, E., Du, N., Hutchinson, B., Pope, R., Bradbury, J., Austin, J., Isard, M., Gur-Ari, G., Yin, P., Duke, T., Levskaya, A., Ghemawat, S., Dev, S., Michalewski, H., Garcia, X., Misra, V., Robinson, K., Fedus, L., Zhou, D., Ippolito, D., Luan, D., Lim, H., Zoph, B., Spiridonov, A., Sepassi, R., Dohan, D., Agrawal, S., Omernick, M., Dai, A. M., Pillai, T. S., Pellat, M., Lewkowycz, A., Moreira, E., Child, R., Polozov, O., Lee, K., Zhou, Z., Wang, X., Saeta, B., Diaz, M., Firat, O., Catasta, M., Wei, J., Meier-Hellstern, K., Eck, D., Dean, J., Petrov, S., Fiedel, N. (2022). *PaLM: Scaling language modeling with Pathways*, arXiv preprint arXiv:2204.02311.
32. Chung, H. W., Hou, L., Longpre, S., Zoph, B., Tay, Y., Fedus, W., Li, Y., Wang, X., Dehghani, M., Brahma, S., Webson, A., Gu, S. S., Dai, Z., Suzgun, M., Chen, X., Chowdhery, A., Castro-Ros, A., Pellat, M., Robinson, K., Valter, D., Narang, S., Mishra, G., Yu, A., Zhao, V., Huang, Y., Dai, A., Yu, H., Petrov, S., Chi, E. H., Dean, J., Devlin, J., Roberts, A., Zhou, D., Le, Q. V., Wei, J. (2024). *Scaling instruction-finetuned language models*, Journal of Machine Learning Research, 25, 1–53, <https://arxiv.org/pdf/2210.11416> (dostęp: 20.07.2025)

33. CONSORT (2025) *CONSORT 2025 statement: updated guidelines for reporting randomized trials*. EQUATOR Network, <https://www.equator-network.org/reporting-guidelines/consort/>, (dostęp: 13.07.2025).
34. COPE (2023). *Authorship and AI tools*. Committee on Publication Ethics, <https://publicationethics.org/guidance/cope-position/authorship-and-ai-tools> (dostęp: 07.07.2025).
35. Council of Europe (2024) *Framework Convention on Artificial Intelligence and Human Rights, Democracy and the Rule of Law*, Council of Europe, <https://www.coe.int/en/web/artificial-intelligence/the-framework-convention-on-artificial-intelligence> (dostęp: 13.07.2025).
36. Creative Commons, <https://creativecommons.org/2025/06/25/introducing-cc-signals-a-new-social-contract-for-the-age-of-ai/>, (dostęp: 07.07.2025).
37. Creswell, J.W., Creswell, J.D. (2018). *Research Design: Qualitative, Quantitative, and Mixed Methods Approaches*, (5th ed.), Thousand Oaks, CA: Sage.
38. Dao, T., Fu, D. Y., Ermon, S., Rudra, A., Ré, C. (2022). *FlashAttention: Fast and memoryefficient exact attention with IO-awareness*, Advances in Neural Information Processing Systems, <https://arxiv.org/pdf/2205.14135> (dostęp: 12.07.2025).
39. Day, T. (2023). *A preliminary investigation of fake peerreviewed citations and references generated by ChatGPT*, The Professional Geographer, 75(6), 1024–1027, <https://doi.org/10.1080/00330124.2023.2190373> (dostęp: 20.07.2025).
40. Dehouche, N. (2021). *Plagiarism in the age of massive Generative Pretrained Transformers (GPT3)*, Ethics in Science and Environmental Politics, 21, 17–23. <https://doi.org/10.3354/esep00195> (dostęp: 10.08.2025).
41. Dettmers, T., Pagnoni, A., Holtzman, A., Zettlemoyer, L. (2023). *QLoRA: Efficient finetuning of quantized LLMs*. Advances in Neural Information Processing Systems, <https://arxiv.org/pdf/2305.14314> (dostęp: 07.07.2025).
42. Ding, J., Ma, S., Dong, L., Zhang, X., Huang, S., Wang, W., Zheng, N., Wei, F. (2023). *LongNet: Scaling Transformers to 1,000,000,000 tokens*, arXiv preprint arXiv:2307.02486, <https://arxiv.org/pdf/2307.02486> (dostęp: 2.08.2025).
43. Dornis, T. W. (2024). *The training of generative AI is not text and data mining*, SSRN Working Paper. <https://ssrn.com/abstract=4993782>, (dostęp: 5.07.2025).
44. dos Santos, A.I. (2023) *The digital competence of academics in higher education*, Education Sciences, 13(2), 132, <https://pmc.ncbi.nlm.nih.gov/articles/PMC9894668/> (dostęp: 13.07.2025).

45. Elsevier. (n.d.), <https://www.elsevier.com/about/policies-and-standards/generative-ai-policies-for-journals> (dostęp: 04.2025).
46. Elsevier. (n.d.), <https://www.elsevier.com/about/policies-and-standards/publishing-ethics> (dostęp: 04.2025).
47. Elsevier. (n.d.), <https://www.elsevier.com/about/policies-and-standards/publishing-ethics-books/the-use-of-generative-ai-and-ai-assisted-technologies-in-the-editing-process> (dostęp: 04.2025).
48. Emerald Publishing. (n.d.), <https://www.emeraldgrouppublishing.com/news-and-press-releases/emerald-publishings-stance-ai-tools-content-creation-and-peer-review> (dostęp: 04.2025).
49. Emerald Publishing. (n.d.), <https://www.emeraldgrouppublishing.com/publish-with-us/ethics-integrity/research-publishing-ethics> (dostęp: 04.2025).
50. Fagbohun, O., Harrison, R.M. i Dereventsov, A. (2024). *An Empirical Categorization of Prompting Techniques for Large Language Models: A Practitioner's Guide*, <https://arxiv.org/abs/2402.14837> (dostęp: 07.07.2025).
51. Fan, W., Ding, Y., Ning, L., Wang, S., Li, H., Yin, D., Chua, T.-S., Li, Q. (2024). *A Survey on RAG Meeting LLMs: Towards RetrievalAugmented Large Language Models*. Proceedings of the 30th ACM SIGKDD Conference on Knowledge Discovery and Data Mining, 6491–6501, <https://doi.org/10.1145/3637528.3671470> (dostęp: 04.08.2025).
52. Farquhar, S., Cuccu, G., Sayers, A., i in. (2024). *Detecting hallucinations in large language models using statistics*, Nature, <https://doi.org/10.1038/s41586-024-07421-0>
53. Feretzakis, G., Chondrogiorgi, M., Gkoulalas-Divanis, A., Verykios, VS. (2024). *Privacy preserving techniques in generative AI and LLMs: A review*, Information, 15(11), 697, <https://doi.org/10.3390/info15110697> (dostęp: 08.08.2025).
54. Ferrara, E. (2023). *Should ChatGPT be biased? Challenges and risks of bias in large language models*, First Monday, 28(11) <https://arxiv.org/abs/2304.03738> (dostęp 10.07.2025).
55. Feuerriegel, S., Hartmann, J., Janiesch, C., Zschech, P. (2024). *Generative AI. Business & Information Systems Engineering*, 66(1), 111–126, <https://doi.org/10.1007/s12599-023-00834-7> (dostęp: 07,07,2025).
56. Fitria, T., N. (2021), “Grammarly” as AIpowered English writing assistant: Students’ alternative for English writing. Metathesis: Journal of English Language, Literature, and Teaching, 5(1), 65–78. <https://doi.org/10.31002/metathesis.v5i1.3519> (dostęp: 07.07.2025).

57. Flanagan, A., BibbinsDomingo, K., Berkwits, M., Christiansen, S. L. (2023). *Nonhuman “authors” and implications for the integrity of scientific publication and medical knowledge*, JAMA, 329(8), 637–639.
58. Friesike, S. (2014). *Open Science: One Term, Five Schools of Thought*. In: S. Bartling, S. Friesike (eds), *Opening Science*, Springer, https://link.springer.com/chapter/10.1007/978-3-319-00026-8_2 (dostęp: 07.07.2025).
59. Gaebel, M., Zhang, T., Stoeber, H. and Morrisroe, A. (2023) *The future of digitally enhanced learning and teaching in European higher education institutions: Final report of the DIGI-HE project*, Brussels: European University Association, <https://eua.eu/resources/publications/1046:digi-he-final-report.html> (dostęp: 13.07.2025).
60. Ganjavi, C., Eppler, M. B., Pekcan, A., Biedermann, B., Abreu, A., Collins, G. S., Gill, I. S., Cacciamani, G., E. (2024). *Publishers’ and journals’ instructions to authors on use of generative AI in academic and scientific publishing: Bibliometric analysis*, BMJ, 384, e077192. <https://doi.org/10.1136/bmj-2023-077192> (dostęp: 07.07.2025).
61. Gao, Y., i in., (2023). *RetrievalAugmented Generation for Large Language Models: A Survey*, arXiv:2312.10997, <https://arxiv.org/abs/2312.10997> (dostęp: 06.07.2025).
62. Gehrmann, S., Strobel, H., Rush, A. M. (2019). *GLTR: Statistical detection and visualization of generated text. Proceedings of the 57th Annual Meeting of the Association for Computational Linguistics: System Demonstrations* (pp. 111–116), Association for Computational Linguistics, <https://aclanthology.org/P19-3019/> (dostęp: 02.08.2025).
63. GodwinJones, R. (2022). *Partnering with AI: Intelligent writing assistance and instructed language learning*. *Language Learning & Technology*, 26(2), 5–24, <https://doi.org/10.10125/73474> (dostęp: 04.08.2025).
64. Google (2023) *Gemini terms of service*, <https://policies.google.com/terms> (dostęp: 07.07.2025).
65. Google (2023) *Prompting Essentials*, Coursera, <https://www.coursera.org/learn/google-prompting-essentials> (dostęp: 07.07.2025).
66. Google Cloud (2024), *What is prompt engineering?* Google Cloud, <https://cloud.google.com/discover/what-is-prompt-engineering>, (dostęp: 07.07.2025).
67. Gravel, J., D’AmoursGravel, M., Osmanliu, E. (2023). *Learning to fake it: Limited responses and fabricated references provided by ChatGPT for medical questions*, Mayo Clinic Proceedings: Digital Health, 1(3), 226–234, <https://doi.org/10.1016/j.mcpdig.2023.05.004>

68. Gustilo, L., Ong, E., Lapinid, M. R. C. (2024). *Algorithmically driven writing and academic integrity: Exploring educators' practices, perceptions, and policies in the AI era*, International Journal for Educational Integrity, 20, Article 3, <https://doi.org/10.1007/s40979-024-00153-8> (dostęp: 02.08.2025).
69. Haber, E., Jemielniak, D., Kurasiński, A., Przegalińska, A. (2025). *Using AI in Academic Writing and Research: A Complete Guide to Effective and Ethical Academic AI*, Springer.
70. He, P., Gao, J., Chen, W. (2021). DeBERTaV3: *Improving DeBERTa using ELECTRA style pretraining with gradient disentangled embedding sharing*, arXiv preprint arXiv:2111.09543, <https://arxiv.org/pdf/2111.09543> (dostęp: 08.07.2025).
71. He, R., Cao, J., Tan, T. (2025). *Generative artificial intelligence: a historical perspective*. National Science Review, 12(5), nwaf050, <https://doi.org/10.1093/nsr/nwaf050> (dostęp: 08.07.2025).
72. Hemminki-Reijonen, U., Hassan, N. M. A. M., Huotilainen, M., Koivisto, J.-M., Cowley, B. U. (2025). *Design of generative AI-powered pedagogy for virtual reality environments in higher education*, npj Science of Learning, 10, Article 31, <https://www.nature.com/articles/s41539-025-00326-1> (dostęp: 02.08.2025).
73. Ho, J., Jain, A., Abbeel, P. (2020). *Denoising Diffusion Probabilistic Models*, arXiv:2006.11239, <https://doi.org/10.48550/arXiv.2006.11239> (dostęp: 07.07.2025).
74. Hoffmann, J., (2022). *Training Compute-Optimal Large Language Models*. *Advances in Neural Information Processing Systems*, 35, NeurIPS 2022, Version proceedings.
75. Hoffmann, J., Borgeaud, S., Mensch, A., Buchatskaya, E., Cai, T., Rutherford, E., ... Sifre, L. (2022). *Training compute-optimal large language models*. *Advances in Neural Information Processing Systems*, <https://arxiv.org/abs/2203.15556> (dostęp: 02.08.2025).
76. Hosseini, M., Resnik, D. B., Holmes, K. (2023). *The ethics of disclosing the use of artificial intelligence tools in writing scholarly manuscripts*, Research Ethics, 19(4), 449–465, <https://doi.org/10.1177/17470161231180449> (dostęp: 12.06.2025).
77. Hu, E. J., Shen, Y., Wallis, P., AllenZhu, Z., Li, Y., Wang, S., Wang, L., Chen, W. (2022). *LoRA: Lowrank adaptation of large language models*, International Conference on Learning Representations.
78. IEEE (n.d.), <https://journals.ieeeauthorcenter.ieee.org/become-an-ieee-journal-author/publishing-ethics/guidelines-and-policies/submission-and-peer-review-policies/> (dostęp: 04.2025)
79. IEEE (n.d.), <https://open.ieee.org/author-guidelines-for-artificial-intelligence-ai-generated-text/> (dostęp: 04.2025)

80. Imran (2024), *Multimodal composing with generative AI: Examining preservice teachers' processes and perspectives*, <https://experts.azregents.edu/en/publications/multimodal-composing-with-generative-ai-examining-preservice-teac> (dostęp: 10.08.2025)
81. Inagaki, T., Kato, A., Takahashi, K., Ozaki, H., Kanda, G. N. (2023). *LLMs can generate robotic scripts from goaloriented instructions in biological laboratory automation*, arXiv:2304.10267, <https://arxiv.org/abs/2304.10267> (dostęp: 02.08.2025).
82. Jakesch, M., Bhat, A., Buschek, D., Zalmanson, L., Naaman, M. (2023). *CoWriting with Opinionated Language Models Affects Users' Views*, In CHI '23, <https://doi.org/10.1145/3544548.3581196> (dostęp: 02.08.2025).
83. Jiang, T., i in. (2024). *A usercentered research agenda for humanAI interaction*, Patterns, 5(5), 100955. <https://doi.org/10.1016/j.patter.2024.100955> (dostęp: 02.08.2025).
84. Journal of Modern Science - Polityka AI, Journal of Modern Science, <https://www.jomswsge.com/PolitykaAI,5588.html> (dostęp: 10.08.2025).
85. Jose, B., Cherian, J., Verghis, A.M., Varghise, S.M., Mumthas, S. i Joseph, S. (2025) *The cognitive paradox of AI in education: between enhancement and erosion*, Frontiers in Psychology, 16:1550621, <https://www.frontiersin.org/journals/psychology/articles/10.3389/fpsyg.2025.1550621/ful>, (dostęp: 04.07.2025).
86. Kafai, Y. & Resnick, M. (eds.), 1996. *Constructionism in Practice: Designing, Thinking, and Learning in a Digital World*, Mahwah, NJ: Lawrence Erlbaum Associates.
87. Khalifa, M., Albadawy, M. (2024). *Using artificial intelligence in academic writing and research: An essential productivity tool*, Computer Methods and Programs in Biomedicine Update, 5(1), 100145, <https://doi.org/10.1016/j.cmpbup.2024.100145> (dostęp: 08.08.2025).
88. Kirchenbauer, J., Geiping, J., Wen, Y., Katz, J., Miers, I., & Goldstein, T. (2023). *A watermark for large language models*. W: *Proceedings of the 40th International Conference on Machine Learning*, PMLR, Vol. 202, pp. 17061–17084.
89. Kirchner, J. H., Ahmad, L., Aaronson, S., Leike, J. (2023). *New AI classifier for indicating AI-written text*, OpenAI, <https://openai.com/index/new-ai-classifier-for-indicating-ai-written-text/> (dostęp: 02.08.2025).
90. Kolb, A., Kolb D., (2022) *Uczenie na podstawie doświadczenia. Podręcznik dla trenerów, edukatorów, coachów*, Warszawa: Dialogi & Zmysły.
91. Krishna, K., Song, Y., Karpinska, M., Wieting, J. F., & Iyyer, M. (2023). *Paraphrasing evades detectors of AI-generated text, but retrieval is an effective defense*. arXiv preprint arXiv:2303.13408 (dostęp: 10.07.2025)

92. Kuru, T. (2024). *Lawfulness of the mass processing of publicly accessible online data to train large language models*, *International Data Privacy Law*, 14(4), 326–339.
<https://doi.org/10.1093/idpl/ipae013> (dostęp: 02.08.2025).
93. Lampropoulos, G. (2025). *Combining Artificial Intelligence with Augmented Reality and Virtual Reality in Education: Current Trends and Future Perspectives*, *Multimodal Technologies and Interaction*, 9(2), 11, <https://www.mdpi.com/2414-4088/9/2/11>, (dostęp: 10.08.2025).
94. Lee, M.K., Microsoft Research, Carnegie Mellon University (2025) *AI and Critical Thinking: Survey Results and Implications*, Redmond, WA: Microsoft Research, Dostęp: https://www.microsoft.com/en-us/research/wp-content/uploads/2025/01/lee_2025_ai_critical_thinking_survey.pdf (dostęp: 07.07.2025).
95. Levy, N. (2025). *Responsibility is not required for authorship*. *Journal of Medical Ethics*. Advance online publication, <https://pmc.ncbi.nlm.nih.gov/articles/PMC12015057/> (dostęp: 07.07.2025).
96. Li, H., Fang, Y., Zhang, S., Lee, S. M., Wang, Y., Trexler, M., Botelho, A. F. (2025). *ARCHED: A Human-Centered Framework for Transparent, Responsible, and Collaborative AI-Assisted Instructional Design*, <https://arxiv.org/abs/2503.08931> (dostęp: 10.08.2025).
97. Li, K. (2024). *Copyright protection during the training stage of generative AI*, *Computer Law & Security Review*. <https://doi.org/10.1016/j.clsr.2024> (dostęp: 02.08.2025).
98. Li, X.L. i Liang, P. (2021) *Prefix-Tuning: Optimizing Continuous Prompts for Generation*.
99. Li, Z., Liang, C., Peng, J., & Yin, M. (2024). *The value, benefits, and concerns of generative ai-powered assistance in writing*, In *Proceedings of the 2024 CHI Conference on Human Factors in Computing Systems* (pp. 1-25).
100. Liang, W., Yuksekgonul, M., Mao, Y., Wu, E., Zou, J. (2023). *GPT detectors are biased against non-native English writers*, *Patterns*, 4(7), 100779,
<https://www.sciencedirect.com/science/article/pii/S2666389923001307> (dostęp: 12.08.2025).
101. Liu, P., Yuan, W., Fu, J., Jiang, Z., Hayashi, H. i Neubig, G. (2023) *Pre-train, Prompt, and Predict: A Systematic Survey of Prompting Methods in Natural Language Processing*, *ACM Computing Surveys*, 55(9), s. 1–35, <https://dl.acm.org/doi/10.1145/3560815> (dostęp: 07.07.2025).
102. Liu, X., Ji, J., Fu, Y., Tam, W.L., Tang, J. i Han, J. (2021) *P-tuning v2: Prompt tuning can be comparable to fine-tuning universally across scales and tasks*, <https://aclanthology.org/2022.acl-short.8> (dostęp: 07.07.2025).

103. Lu, Z., Peng, Y., Cohen, T., Ghassemi, M., Weng, C., Tian, S. (2024). *Large language models in biomedicine and health: current research landscape and future directions*, Journal of the American Medical Informatics Association, 31(9), 1801–1811, <https://doi.org/10.1093/jamia/ocae202> (dostęp: 08.08.2025).
104. Marton, F., Säljö, R. (1976) *On qualitative differences in learning: I -Outcome and process*, British Journal of Educational Psychology, 46(1), 4–11. (dostęp: 06.06.2025).
105. Microsoft (2023) *Microsoft announces new Copilot Copyright Commitment for customers* <https://blogs.microsoft.com/on-the-issues/2023/09/07/copilot-copyright-commitment-ai-legal-concerns/> (dostęp: 09.08.2025).
106. Microsoft (2024). *Learn about Copilot prompts. Microsoft Support*, <https://support.microsoft.com/en-us/topic/learn-about-copilot-prompts-f6c3b467-f07c-4db1-ae54-ffac96184dd5> (dostęp: 07.07.2025).
107. Microsoft. (2025). *Data, privacy, and security for Microsoft 365 Copilot, Microsoft Learn*, <https://learn.microsoft.com/en-us/copilot/microsoft-365/microsoft-365-copilot-privacy> (dostęp: 02.08.2025).
108. Minaee, S., Mikolov, T., Nikzad, N., Chenaghlu, M., Socher, R., Amatriain, X., Gao, J. (2024). *Large language models: A survey*, arXiv preprint arXiv:2402.06196, <https://doi.org/10.48550/arXiv:2402.06196> (dostęp: 02.08.2025).
109. Mitchell, E., Lee, Y., Khazatsky, A., Manning, C. D., Finn, C. (2023). *DetectGPT: Zero-shot machine-generated text detection using probability curvature*, arXiv preprint arXiv:2301.11305, <https://arxiv.org/abs/2301.11305> (dostęp: 07.07.2025).
110. Mittelstadt, B., Russell, C., Wachter, S. (2023). *To protect science we must use LLMs as zero-shot translators*, Nature Human Behaviour, 7(12), 1645–1648, <https://doi.org/10.1038/s41562-023-01535-5> (dostęp: 06.06.2025).
111. Mosqueira-Rey, E., Hernández-Pereira, E., Alonso-Ríos, D., Bobes-Bascarán, J., Fernández-Leal, Á. (2023). *Human-in-the-loop machine learning: a state of the art*, Artificial Intelligence Review, 56(4), 3005–3054, <https://doi.org/10.1007/s10462-022-10246-w> (dostęp: 06.06.2025).
112. Mueller, F. B., Gorge, R., Bernzen, A. K., Pirk, J. C., & Poretschkin, M. (2024). *LLMs and memorization: On quality and specificity of copyright compliance*, In Proceedings of the AAAI/ACM Conference on AI, Ethics, and Society (AIES) (pp. 984–993). <https://doi.org/10.1609/aies.v7i1.31697> (dostęp: 02.06.2025).
113. Mytnik, J. (2025). *AI dla edukacji: narzędzia i metaprompty* [wideo], Centrum Nowoczesnej Edukacji, Politechnika Gdańska, YouTube, <https://youtu.be/LiiNulqQV68> (dostęp: 10.08.2025).

114. Mytnik, J. (2025) *Iluzja mistrzostwa: GenAI i uczenie się*, [wideo] YouTube, <https://youtu.be/mShQ5ngrlHc>, (dostęp: 10.08.2025).
115. National Institute of Standards and Technology (2023) *Artificial Intelligence Risk Management Framework (AI RMF 1.0)*, Gaithersburg, MD: U.S. Department of Commerce. <https://nvlpubs.nist.gov/nistpubs/ai/nist.ai.100-1.pdf>, (dostęp: 13.07.2025).
116. National Research Council (2015) *Enhancing the effectiveness of team science*, Washington, DC: The National Academies Press, <https://nap.nationalacademies.org/read/19007/chapter/1#x> (dostęp: 13.07.2025).
117. Nature Portfolio (2023) *Artificial intelligence (AI) – editorial policies*, Nature, <https://www.nature.com/nature-portfolio/editorial-policies/ai> (dostęp: 13.07.2025).
118. Novelli, C., De Angelis, A., Gianfrancesco, E., Floridi, L. (2024). *Generative AI in EU law: Liability, privacy, intellectual property, and cybersecurity*, Computer Law & Security Review (preprint version), arXiv:2401.07348, <https://arxiv.org/abs/2401.07348>, (dostęp: 07.07.2025).
119. Nowotny, H. (2000). *Transgressive competence: The narrative of expertise*, European Journal of Social Theory, 3(1), 5–21, <https://journals.sagepub.com/doi/10.1177/136843100003001001> (dostęp: 12.07.2025).
120. OpenAI (2023). *Prompt engineering best practices for ChatGPT*, OpenAI Help Center, <https://help.openai.com/en/articles/10032626-prompt-engineering-best-practices-for-chatgpt> (dostęp: 05.08.2025).
121. OpenAI (2025) *Europe Terms of use*, <https://openai.com/policies/terms-of-use> (dostęp: 07.07.2025).
122. Originality.ai, *AI Detector - Most Accurate AI Content Checker for ChatGPT & Gemini*, <https://originality.ai/> (dostęp: 05.08.2025).
123. OrduñaMalea, E., CabezasClavijo, Á. (2023). *ChatGPT and the potential growing of ghost bibliographic references*, Scientometrics, 128(9), 5351–5355, <https://doi.org/10.1007/s11192-023-04804-4> (dostęp: 07.07.2025).
124. Ouyang, L., Wu, J., Jiang, X., Almeida, D., Wainwright, C., Mishkin, P., Zhang, C., Agarwal, S., Slama, K., Ray, A., Schulman, J., Hilton, J., Kelton, F., Miller, L., Simens, M., Aspell, A., Welinder, P., Christiano, P. F., Leike, J., Lowe, R. (2022). *Training language models to follow instructions with human feedback*. *Advances in Neural Information Processing Systems*, <https://arxiv.org/pdf/2203.02155> (dostęp: 12.08.2025).
125. Owczarczuk, M. (2023). *Ethical and regulatory challenges amid artificial intelligence development: an outline of the issue*, *Ekonomia i Prawo*. Economics and Law, 22(2), 295–310, <https://doi.org/10.12775/EiP.2023.017> (dostęp: 12.07.2025).

126. Page, M.J., McKenzie, J.E., Bossuyt, P.M., Boutron, I., Hoffmann, T.C., Mulrow, C.D., Shamseer, L., Tetzlaff, J.M., Akl, E.A., Brennan, S.E., Chou, R., Glanville, J., Grimshaw, J.M., Hróbjartsson, A., Lalu, M.M., Li, T., Loder, E.W., Mayo-Wilson, E., McDonald, S., McGuinness, L.A., Stewart, L.A., Thomas, J., Tricco, A.C., Welch, V.A., Whiting, P. and Moher, D. (2021). *The PRISMA 2020 statement: An updated guideline for reporting systematic reviews*, BMJ, 372, n71, <https://www.bmj.com/content/372/bmj.n71> (dostęp: 13.07.2025).
127. Papert, S. (1996) *Burze mózgów, Dzieci i komputery*, Wydawnictwo PWN.
128. Partnership for 21st Century Skills (2003) *Learning for the 21st Century: A Report and MILE Guide for 21st Century Skills*, Washington, D.C.
129. Peng, S., Kalliamvakou, E., Cihon, P., Demirer, M. (2023), *The Impact of AI on Developer Productivity: Evidence from GitHub Copilot*, <https://arxiv.org/abs/2302.06590> (dostęp: 06.06.2025).
130. Piwowar, H., Priem, J., Orr, R. (2019). The state of OA: *A large-scale analysis of Open Access adoption*, PeerJ, <https://peerj.com/articles/4375/>, (dostęp: 07.07.2025).
131. Polska Akademia Nauk (2024), *Kodeks etyki pracownika naukowego*, Warszawa: PAN https://pan.pl/app/uploads/2025/06/Kodeks_etyki_pracownika_naukowego_2024.pdf, (dostęp: 07.07.2025).
132. Politechnika Gdańska, Centrum Nowoczesnej Edukacji (2024), *Baza inspiracji GenAI*, <https://cne.pg.edu.pl/genai/baza-inspiracji>, (dostęp: 10.08.2025).
133. Press, O., Smith, N. A., Lewis, M. (2022). *Train short, test long: Attention with linear biases enables input length extrapolation*, International Conference on Learning Representations.
134. Przeglasińska, A., Jemielniak, D. (2023). *AI w strategii marketingowej*, MT Biznes.
135. Przeglasińska, A., Triantoro, T. (2024). *Przenikanie umysłów. Potencjał twórczy współpracy z AI*, AI Books by CampusAI.
136. Puentedura, R.R. (2013). *SAMR: Getting To Transformation*, Hippasus Weblog, <http://www.hippasus.com/rrpweblog/archives/2013/04/16/SAMRGettingToTransformation.pdf> (dostęp: 07.07.2025).
137. QuillBot. (2025). *AI Content Detector*. <https://quillbot.com/ai-content-detector> (dostęp: 10.08.2025).
138. Quintais, J. P. (2025). *Generative AI, copyright and the AI Act*, Computer Law & Security Review. <https://doi.org/10.1016/j.clsr.2025> (dostęp: 12.08.2025).

139. Rafailov, R., Sharma, A., Mitchell, E., Ermon, S., Manning, C. D., Finn, C. (2023). *Direct Preference Optimization: Your language model is secretly a reward model*, Advances in Neural Information Processing Systems, <https://arxiv.org/pdf/2305.18290> (dostęp: 12.07.2025).
140. Randomized Controlled Trial of Generative AI Coding Tools in the Workplace (2024), <https://arxiv.org/abs/2410.18334> (dostęp: 13.07.2025).
141. Redecker, C. (2017) *European framework for the digital competence of educators: DigCompEdu*, Luxembourg: Publications Office of the European Union, <https://publications.jrc.ec.europa.eu/repository/handle/JRC107466> (dostęp: 13.07.2025).
142. Resnick, M. (2017). *Lifelong kindergarten: Cultivating creativity through projects, passion, peers, and play*. MIT Press.
143. Resnick, M. (2025) *Generative AI and creative learning: Concerns, opportunities, and choices*, Medium, <https://mres.medium.com/ai-and-creative-learning-concerns-opportunities-and-choices-63b27f16d4d0> (dostęp: 07.07.2025).
144. Resnick, M. (2025) *In the age of AI, we need a human-centered society more than ever*, Medium, <https://mres.medium.com/in-the-age-of-ai-we-need-a-human-centered-society-more-than-ever-0be3449ad40e> (dostęp: 07.07.2025).
145. Resnik, D. B., Hosseini, M. (2024). *The ethics of using artificial intelligence in scientific research: New guidance needed for a new tool. AI and Ethics*. Springer.
146. Rezaev, A. V. (2021). *Twelve Theses on Artificial Intelligence and Artificial Sociality*, Monitoring of Public Opinion: Economic and Social Changes, 1, 20–30, <https://doi.org/10.14515/monitoring.2021.1.1894> (dostęp: 13.07.2025).
147. Ruschemeier, H. (2025). *Generative AI and data protection*, Cambridge Forum on AI Law and Governance (accepted manuscript), SSRN 4814999.
148. Sadasivan, V. S., Kumar, A., Balasubramanian, S., Wang, W., Feizi, S. (2023). *Can AI-generated text be reliably detected? Transactions on Machine Learning Research (TMLR)*, (Wersja preprint: arXiv:2303.11156), <https://arxiv.org/abs/2303.11156> (dostęp: 02.06.2025).
149. Sage. (n.d.). <https://www.sagepub.com/about/policies/ai-author-guidelines> (dostęp: 12.04.2025).
150. Sage. (n.d.). <https://www.sagepub.com/journals/editorial-policies/artificial-intelligence-policy> (dostęp: 12.04.2025).
151. Sahoo, P., Singh, A.K., Saha, S., Jain, V., Mondal, S., Chadha, A. (2024). *A Systematic Survey of Prompt Engineering in Large Language Models: Techniques and Applications*, <https://arxiv.org/abs/2402.07927> (dostęp: 07.07.2025).

152. Schick, T., Dwivedi Yu, J., Dessì, R., Raileanu, R., Lomeli, M., Zettlemoyer, L., Cancedda, N., Scialom, T. (2023). *Toolformer: Language Models Can Teach Themselves to Use Tools*, arXiv:2302.04761, <https://arxiv.org/abs/2302.04761> (dostęp: 2.06.2025).
153. Springer Nature (2023) *Editorial policies*, Springer Nature, <https://www.springernature.com/gp/policies/editorial-policies> (dostęp: 13.07.2025).
154. Taylor & Francis. (n.d.), <https://taylorandfrancis.com/about/ai/> (dostęp: 30.04.2025).
155. Taylor & Francis. (n.d.), <https://taylorandfrancis.com/our-policies/ai-policy/> (dostęp: 30.04.2025).
156. Tomczyk, P., Brüggemann, P., Mergner, N., Petrescu, M. (2024a). *Are AI tools better than traditional tools in literature searching? Evidence from E-commerce research*, Journal of Librarianship and Information Science, <https://journals.sagepub.com/doi/10.1177/09610006241295802> (dostęp: 07.07.2025).
157. Tomczyk, P., Brüggemann, P., Mergner, N., Petrescu, M. (2024b). *Exploring AI's role in literature searching: Traditional methods versus AI-based tools in analyzing topical E-commerce themes*. W: F. J. Martínez-López (red.), *Advances in Digital Marketing and eCommerce - Fifth International Conference*, Springer.
158. Tomczyk, P., Brüggemann, P., Paul, J. (2024). *Variable science mapping: A methodological extension for interdisciplinary research*, Journal of Marketing Analytics, <https://doi.org/10.1057/s41270024002474> (dostęp: 02.08.2025).
159. Tomczyk, P., Brüggemann, P., Paul, J. (2024). *Variable science mapping as literature review method*, Journal of Marketing Analytics.
160. Tomczyk, P., Brüggemann, P., Vrontis, D. (2024). *AI meets academia: Transforming systematic literature reviews*, EuroMed Journal of Business.
161. Touvron, H., Lavril, T., Izacard, G., Martinet, X., Lachaux, M.-A., Lacroix, T., Rozière, B., Goyal, N., Hambro, E., Azhar, F., Rodriguez, A., Joulin, A., Grave, E., Lample, G. (2023). *LLaMA: Open and efficient foundation language models*, arXiv preprint arXiv:2302.13971.

162. Touvron, H., Martin, L., Stone, K., Albert, P., Almahairi, Y., Babaei, Y., Bashlykov, N., Batra, S., Bhargava, P., Bhosale, S., Bikel, D., Blecher, L., Canton Ferrer, C., Chen, M., Cucurull, G., Esioibu, D., Fernandes, J., Fu, J., Fu, W., Fuller, B., Gao, C., Goswami, V., Goyal, N., Hartshorn, A., Hosseini, S., Hou, R., Inan, H., Kardaş, M., Kerkez, V., Khabsa, M., Kloumann, I., Korenev, A., Koura, P. S., Lachaux, M.-A., Lavril, T., Lee, J., Liskovich, D., Lu, Y., Mao, Y., Martinet, X., Mihaylov, T., Mishra, P., Molybog, I., Nie, Y., Poulton, A., Reizenstein, J., Rungta, R., Saladi, K., Schelten, A., Silva, R., Smith, E. M., Subramanian, R., Tan, X. E., Tang, B., Taylor, R., Williams, A., Kuan, J. X., Xu, P., Yan, Z., Zarov, I., Zhang, Y., Fan, A., Kambadur, M., Narang, S., Rodriguez, A., Stojnic, R., Edunov, S., Scialom, T. (2023). *Llama 2: Open foundation and finetuned chat models*, arXiv preprint arXiv:2307.09288, <https://arxiv.org/pdf/2307.09288> (dostęp: 08.07.2025).
163. UNESCO (2023). *Guidance for generative AI in education and research*, Paris: United Nations Educational, Scientific and Cultural Organization, <https://www.unesco.org/en/articles/guidance-generative-ai-education-and-research> (dostęp: 02.06.2025).
164. Vuorikari, R., Kluzer, S. and Punie, Y. (2022) *DigComp 2.2: The digital competence framework for citizens – With new examples of knowledge, skills and attitudes*, Luxembourg: Publications Office of the European Union, <https://publications.jrc.ec.europa.eu/repository/handle/JRC128415> (dostęp: 13.07.2025).
165. Wagner, M. W., ErtlWagner, B. B. (2024). *Accuracy of information and references using ChatGPT3 for retrieval of clinical radiological information*, Canadian Association of Radiologists Journal, 75(1), 69–73, <https://doi.org/10.1177/08465371231171125> (dostęp: 07.07.2025).
166. Walters, W. H. (2023). *The effectiveness of software designed to detect AI-generated writing: A comparison of 16 AI text detectors*, Open Information Science, 7(1), 20220158, <https://www.degruyterbrill.com/document/doi/10.1515/opis-2022-0158> (dostęp: 08.08.2025).
167. Walters, W. H., Wilder, E. I. (2023). *Fabrication and errors in the bibliographic citations generated by ChatGPT*, Scientific Reports, 13, 14045, <https://doi.org/10.1038/s41598-023-41032-5> (dostęp: 05.06.2025).
168. Wang, S., Hu, T., Huang, X., Li, Y., Zhang, C., Ning, H., Zhu, R., Li, Z., & Ye, X. (2024). *GPT, large language models (LLMs) and generative artificial intelligence (GAI) models in geospatial science: A systematic review*. International Journal of Digital Earth, 17(1), 2353122, <https://doi.org/10.1080/17538947.2024.2353122> (dostęp: 07.07.2025).

169. Wang, Yizhong, i in. (2022). *Super-Natural Instructions: Generalization via Declarative Instructions on 1600+ NLP Tasks*, <https://aclanthology.org/2022.emnlp-main.340/> (dostęp: 06.08.2025).
170. Wei, J., Bosma, M., Zhao, V. Y., Guu, K., Yu, A. W., Lester, B., Du, N., Dai, A. M., Le, Q. V. (2022). *Finetuned language models are zeroshot learners*. International Conference on Learning Representations, <https://arxiv.org/pdf/2109.01652> (dostęp: 12.07.2025).
171. White, J., Fu, Q., Hays, S., Sandborn, M., Olea, C., Gilbert, H., Elnashar, A., Spencer-Smith, J. i Schmidt, D.C. (2023). *A Prompt Pattern Catalog to Enhance Prompt Engineering with ChatGPT*, <https://arxiv.org/abs/2302.11382> (dostęp: 07.07.2025).
172. Wiley (n.d.), <https://authorservices.wiley.com/ethics-guidelines/index.html> (dostęp: 30.04.2025).
173. Wiley (n.d.). <https://www.wiley.com/en-us/publish/book/ai-guidelines> (dostęp: 30.04.2025).
174. Wilkinson, M.D., Dumontier, M., Aalbersberg, I.J., Appleton, G., Axton, M., Baak, A., Blomberg, N., Boiten, J.-W., da Silva Santos, L.B., Bourne, P.E., Bouwman, J., Brookes, A.J., Clark, T., Crosas, M., Dillo, I., Dumon, O., Edmunds, S., Evelo, C.T., Finkers, R., Gonzalez-Beltran, A., Gray, A.J., Groth, P., Goble, C., Grethe, J.S., Heringa, J., 't Hoen, P.A.C., Hooft, R., Kuhn, T., Kok, R., Kok, J., Lusher, S.J., Martone, M.E., Mons, A., Packer, A.L., Persson, B., Rocca-Serra, P., Roos, M., van Schaik, R., Sansone, S.-A., Schultes, E., Sengstag, T., Slater, T., Strawn, G., Swertz, M.A., Thompson, M., van der Lei, J., van Mulligen, E., Velterop, J., Waagmeester, A., Wittenburg, P., Wolstencroft, K., Zhao, J. and Mons, B. (2016). *The FAIR guiding principles for scientific data management and stewardship*, Scientific Data, 3, 160018, <https://www.nature.com/articles/sdata201618> (dostęp: 13.07.2025).
175. Wojewodzic, K. (red.) (2024) *WydAIniej. Jak sztuczna inteligencja usprawni twoją pracę*. Wrocław, Oficyna Wydawnicza ATUT.
176. Wu, X., Xiao, L., Sun, Y., Zhang, J., Ma, T., He, L. (2022). *A survey of human-in-the-loop for machine learning*, Future Generation Computer Systems, 135, 364-381, <https://doi.org/10.1016/j.future.2022.05.016> (dostęp: 12.08.2025).
177. Xu, Y., Liu, X., Cao, X., Huang, C., Liu, E., Qian, S., Liu, X., Wu, Y., Dong, F., Qiu, C., Qiu, J., Hua, K., Su, W., Wu, J., Xu, H., Han, Y., Fu, C., Yin, Z., Liu, M., Roepman, R., Dietmann, S., Virta, M., Kengara, F., Zhang, Z., Zhang, L., Zhao, T., Dai, J., Yang, J., Lan, L., Luo, M., Liu, Z., An, T., Zhang, B., He, X., Cong, S., Liu, X., Zhang, W., Lewis, J., Tiedje, J., Wang, Q., An, Z., Wang, F., Zhang, L., Huang, T., Lu, C., Cai, Z., Wang, F., Zhang, J. (2021). *Artificial intelligence: A powerful paradigm for scientific research*. Innovation, 2(4), 100179, <https://doi.org/10.1016/j.xinn.2021.100179> (dostęp: 10.08.2025).

178. Yaacoub, A., Da-Rugna, J., & Assaghir, Z. (2025). *Assessing AI-Generated Questions' Alignment with Cognitive Frameworks in Educational Assessment*.
<https://arxiv.org/abs/2504.14232>, dostę: 10.08.2025).
179. Yan, C., Ren, R., Meng, M. H., Wan, L., Ooi, T. Y., Bai, G. (2024). *Exploring ChatGPT App Ecosystem: Distribution, Deployment and Security*, In ASE '24,
<https://arxiv.org/abs/2408.14357> (dostę: 12.06.2025).
180. Yang, L., Zhang, Z., Song, Y., Hong, S., Xu, R., Zhao, Y., Zhang, W., Cui, B., Yang, M.H. (2022). *Diffusion Models: A Comprehensive Survey of Methods and Applications*, ACM Computing Surveys, 56(3), 1–39, <https://doi.org/10.1145/3626235> (dostę: 12.07.2025).
181. Yao, S., Zhao, J., Yu, D., i in. (2023). *ReAct: Synergizing Reasoning and Acting in Language Models*, In ICLR '23, https://openreview.net/forum?id=WE_vluYUL-X (dostę: 07.07.2025).
182. Yao, Shunyu, i in. (2023). *ReAct: Synergizing Reasoning and Acting in Language Models*, Proceedings of ICLR 2023, <https://arxiv.org/abs/2210.03629> (dostę: 12.07.2025).
183. Yoo, J.H. (2025). *Defining the boundaries of AI use in scientific writing: A comparative review of editorial policies*, Journal of Korean Medical Science, 40, e187,
<https://pmc.ncbi.nlm.nih.gov/articles/PMC12170296> (dostę: 12.08.2025).
184. Zhang, Y., Khan, S. A., Mahmud, A., Yang, H., Lavin, A., Levin, M., Frey, J., Dunnmon, J., Evans, J., Bundy, A., Dzeroski, S., Tegner, J., Zenil, H. (2025). *Exploring the role of large language models in the scientific method: from hypothesis to discovery*, npj Artificial Intelligence, 1(1), 14, <https://www.nature.com/articles/s44387-025-00019-5>, (dostę: 06.07.2025).
185. Zhao, W., Zhou, K., Li, J., Tang, T., Wang, X., Hou, Y., Min, Y., Zhang, B., Zhang, J., Dong, Z., Du, Y., Yang, C., Chen, Y., Chen, Z., Jiang, J., Ren, R., Li, Y., Tang, X., Liu, Z., Liu, P., Nie, J., Wen, J. (2023). *A Survey of Large Language Models*, arXiv:2303.18223,
<https://doi.org/10.48550/arXiv:2303.18223> (dostę: 17.07.2025).
186. Zhao, Y., Llorento P., Gómez S., (2021). *Digital competence in higher education research: A systematic literature review*, Computers & Education, 168, 104212,
<https://www.sciencedirect.com/science/article/pii/S0360131521000890> (dostę: 12.08.2025).
187. ZeroGPT. *AI Detector - Trusted AI Checker for ChatGPT, GPT-5 & Gemini*,
<https://www.zerogpt.com/> (dostę: 08.2025).